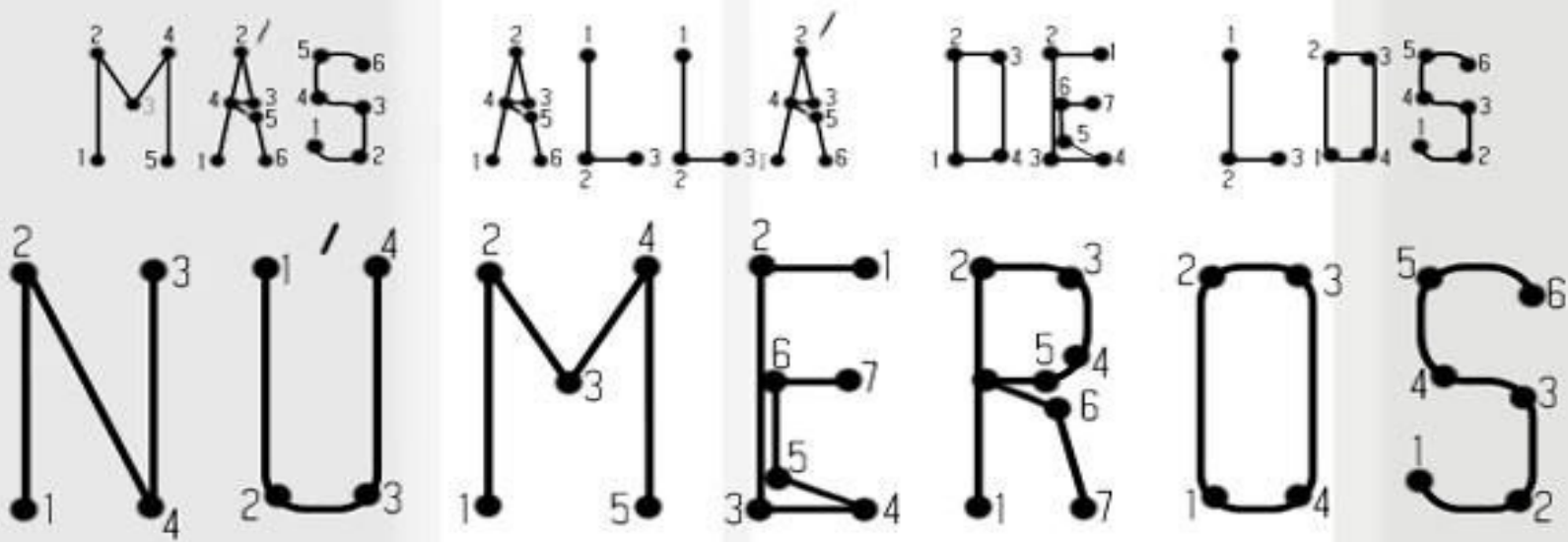




Meditaciones de un matemático

JOHN ALLEN PAULOS



Tal vez sea sintomático de nuestro analfabetismo matemático el hecho de que, en un mundo a todas luces regido por la ciencia y la técnica, estemos cada vez más pendientes de horóscopos y videntes, que nos atraigan los gurús y los paraísos artificiales de los alucinógenos, o que apelemos con mayor virulencia a distintos fundamentalismos religiosos. En *Más allá de los números*, el conocido científico John Allen Paulos nos ofrece precisamente un eficaz antídoto contra esta fobia matemática, una fobia que tal vez proceda de la imagen errónea y confusa que el hombre actual se ha forjado de las matemáticas. Pues «a menudo», escribe Paulos, «ideas matemáticas muy “avanzadas”, son más intuitivas y comprensibles que ciertos temas de álgebra mental».

Concebidos a modo de diccionario, los breves ensayos que constituyen *Más allá de los números* nos invitan a disfrutar de los conceptos de la matemática moderna, a viajar por el interior de la mente que piensa numéricamente. Resultado: una nueva manera de ver el mundo, un nuevo modo de comprender los más triviales sucesos de la vida diaria. Así, el lector verá la utilidad de las matemáticas en, por ejemplo, la composición de música electrónica o en los escrutinios electorales, para resolver problemas económicos caseros, para empapelar una habitación o para entender qué es la inteligencia artificial.



John Allen Paulos

Más allá de los números

Meditaciones de un matemático

Metatemas - 031

ePub r1.2

koothrapali 09.06.15

Título original: *Beyond Numeracy. Ruminations of a Numbers man*

John Allen Paulos, 1991

Traducción: Josep Llosa

Cubierta: TaliZorah

Editor digital: koothrapali

Escaneado: TaliZorah

Imágenes: Mi gatita preferida

ePub base r1.2



A mis padres, Helen y Peter, origen de mis genes X e Y.

Querría dar las gracias también a Rafe Sagalyn, a Robert Frankel, y
a Sheila, Leah y Daniel Paulos.

Introducción

Este libro es en parte un diccionario, en parte una recopilación de ensayos matemáticos cortos y en parte las reflexiones de un matemático. A pesar de contener muchas entradas (ensayos breves) ordenadas alfabéticamente que describen una amplia gama de temas matemáticos, lo que le distingue de un diccionario es que las entradas son menos globales, más largas y, en muchos casos, muy poco convencionales.

Por necesidad, el libro contiene más información que la mayoría de recopilaciones de ensayos. Sin embargo, he intentado mantener el tono personal y unificador típico de éstas. En otras palabras, este libro ha sido escrito por un individuo con sus intereses concretos (no todos matemáticos), sus predisposiciones (las matemáticas como arte liberal y no sólo como herramienta técnica) y sus estrategias pedagógicas (como el empleo de cuentos y aplicaciones poco usuales). Aunque el tema no sea yo, sino las matemáticas, no he hecho ningún esfuerzo por no aparecer en el cuadro, con la esperanza de servir de guía personal al lector a través de un tema que amedrenta a muchos. El público al que me dirijo es inteligente y culto, pero generalmente anumérico (matemáticamente analfabeto).

He recibido una cantidad sorprendente de cartas de lectores de mi anterior libro, *El hombre anumérico*^[1] en las que me manifiestan que éste ha estimulado su interés por las matemáticas y que ahora quieren algo más para satisfacer su recién despertado apetito por el tema —algo del mismo estilo, pero que vaya más allá del simple numerismo—. Cito un tanto impudicamente de la carta de una lectora: «Quizá suene anumérico, pero me gustaría que escribiera otro libro que fuera exactamente igual, sólo que

distinto, algo que avanzara un poco más». Espero que este libro le resulte atractivo y útil, y, al mismo tiempo, consiga ofrecer a sus lectores una imagen no técnica, sin dejar de ser riguroso, de la matemática y su relación con el mundo que nos rodea.

Hay muchísimas personas que aprecian la belleza y la importancia de las matemáticas, pero que, como no pueden volver a la universidad, no ven la manera de profundizar en este interés. Algo les ha llevado a creer que sin ningún conocimiento de los formalismos, teoremas y manipulaciones simbólicas, las ideas matemáticas están por completo fuera de su alcance. Creo que esto es falso y, lo que es peor, totalmente pernicioso. Se puede aprender de Montaigne, Flaubert y Camus sin saber leer en francés, y del mismo modo se puede aprender de Euler, Gauss y Gödel sin resolver ecuaciones diferenciales. Lo que hace falta en ambos casos es un traductor que maneje bien ambos idiomas.

Como aspirante a esa especie de traductor, he procurado evitar, en la medida de lo posible, ecuaciones, tablas y diagramas complicados, así como símbolos formales. Incluyo unas cuantas ilustraciones y hago breves menciones de algunas notaciones matemáticas corrientes, porque resultan a veces indispensables y son especialmente útiles si se consultan otros libros. Sin embargo, la mayor parte de lo expuesto se hace con palabras, y en nuestra lengua de comunicación.

Las entradas van desde los resúmenes de disciplinas enteras (cálculo, trigonometría, topología) a notas biográficas e históricas (Gödel, Pitágoras, geometría no euclídea), pasando por fragmentos de folklore matemático o cuasimatemático (conjuntos infinitos, poliedros regulares, QED) muy conocidos por los matemáticos pero no por los profanos, aunque sean personas cultas. He incluido de vez en cuando fragmentos menos convencionales: la reseña de un libro inexistente, un «fluir matemático de conciencia» durante un viaje en coche y breves discusiones sobre humor o ética. Se han tratado temas nuevos (el caos y los fractales, la iteración, la complejidad) y también otros más clásicos (las secciones cónicas, la inducción matemática, los números primos).

Soy totalmente culpable de cometer flagrantes «errores de categoría» a lo largo de toda la obra: al incluir como entradas temas matemáticos, principios

pedagógicos, pequeñas homilías y anécdotas, como si todo ello estuviera coordinado. No pido disculpas por ello, pues estas discusiones tan dispares ilustran un hecho que a menudo se pasa por alto: que la matemática es una empresa humana con muchos estratos, y no simplemente un conjunto de teoremas y cálculos formales.

Escribir artículos matemáticos no es lo mismo que escribir sobre la matemática, pero pienso que no tendría por qué haber un abismo tan grande entre ambas actividades (a menudo he soñado con anunciar la solución de un problema famoso en un libro de divulgación en vez de en una revista especializada tradicional). En lo que respecta a la precisión de las diversas entradas, he intentado seguir un rumbo difícil de mantener: escribir con la precisión suficiente para evitar el desdén académico (el desinterés académico por este tipo de obras de divulgación es inevitable) y, sin embargo, con la claridad suficiente para evitar que los lectores se formen conceptos falsos. Cuando la claridad y la precisión están en conflicto, como ocurre a veces, he optado la mayoría de las veces por la primera.

Una idea errónea muy extendida es que la matemática es completamente jerárquica: primero la aritmética, luego el álgebra, después el cálculo, a continuación más abstracción y luego lo que sea. (¿Qué viene después del cálculo superior? Respuesta: una parálisis grave). Esta creencia en la condición de poste totémico de la matemática es falsa, pero lo peor es que impide que muchas personas que pasaron apuros para aprobar las matemáticas en la enseñanza básica, en la escuela secundaria o incluso en la universidad tomen un libro de divulgación sobre el tema. A menudo, ideas matemáticas muy «avanzadas» son más intuitivas y comprensibles que ciertos temas de álgebra elemental. Mi lema es: si te quedas atascado y no entiendes algo, sigue adelante y probablemente la niebla se levantará, a menudo antes de acabar el artículo.

Para acabar, recuerdo a las personas que he conocido que, teniéndose por anuméricas, se han sorprendido al comprobar su intuición matemática. Al tener una idea calculística tradicional de la matemática, esas personas suelen caracterizar sus comentarios perspicaces como lógica o sentido común, nunca como matemática; me recuerdan al burgués de Molière que se sorprendió al descubrir que llevaba toda la vida hablando en prosa. Este libro está pues

escrito para los matematófilos que no saben que lo son (entre otros), que toda la vida han pensado matemáticamente sin haberlo notado.

Las entradas son generalmente independientes y a veces se cruzan referencias.

Al estilo matemático

Aunque casi todo el mundo reconoce la importancia práctica de estudiar matemáticas, relativamente pocos aceptarán que la matemática de la vida cotidiana pueda ser un tema atractivo para la reflexión ociosa. Sin embargo, la matemática proporciona un modo de entender el mundo, y el hecho de desarrollar una conciencia o una perspectiva matemática puede ayudarnos en nuestro comportamiento cotidiano.

En vez de razonar esto último lo ilustraré con una anécdota. Recientemente tuve que desplazarme a Nueva York con una cierta urgencia y llevaba un poco de prisa. Mientras guardaba cola para llegar al peaje iban creciendo en mí los pensamientos asesinos habituales cuando me di cuenta de que el conductor del primer coche de mi fila estaba dejando que otros coches de la fila de su derecha, que estaba muy llena, le (y me) adelantaran. Había un semáforo en el cruce, por tanto no hacía falta dar esas muestras de filantropía, y el aspirante a samaritano debería haber considerado que su buena acción suponía también un perjuicio a los conductores que estaban detrás de él. En este caso, la integral matemática o suma de estos inconvenientes era mayor. Aunque no se trate, ni mucho menos, de una reflexión profunda, este «cálculo» y otros similares parecen totalmente ajenos a muchas personas.

Al llegar por fin a la autopista, aceleré rápidamente hasta alcanzar una velocidad media de unos 110 kilómetros por hora, reduciendo hasta los 80 sólo cuando aparecía algún coche patrulla. A pesar de mi carrera, la necesidad de este juego parecía ese día especialmente clara y me pregunté cómo nunca nadie había puesto en práctica una idea tan simple como la siguiente para reprimir el exceso de velocidad en las autopistas de peaje: cuando alguien entra en una de esas vías recoge un billete con la hora de entrada impresa.

Como se conoce la distancia entre los distintos puestos de peaje, cuando el ordenador imprime la hora de salida se puede calcular fácilmente la velocidad media de dicha persona durante el trayecto. El encargado del peaje podría entonces enviar a los conductores con billetes incriminadores a un coche patrulla estacionado allí mismo.

Este método no acabaría con todos los excesos de velocidad, naturalmente, pues uno podría conducir muy aprisa hasta exactamente antes de la salida, detenerse y tomar una taza de café o hacer una comida completa si hubiera corrido de verdad, y salir con una velocidad media legal. Sin embargo, el aliciente primario del exceso de velocidad se habría eliminado. ¿Qué tiene de malo este plan? Dividir un número por otro, la distancia recorrida por el tiempo empleado, no es seguramente una técnica arriesgada ni novedosa. En la actualidad se ponen multas por exceso de velocidad basándose en el radar, que es mucho menos fiable.

Puse la radio para escapar de estos pensamientos y me acordé de cómo me gustaría, aunque sólo fuera una vez, oír una pieza de rock que usara la palabra *doesn't* en vez de *don't*, como en *She don't love me anymore* («Ella no me quiere») o como la que estaban tocando entonces, *It don't matter anyway*. («De todos modos no importa»).[2] Tal vez por el relativo entumecimiento sensorial de conducir, se me pegó esta triste letra. Quizá no importaba a pesar de todo y, si así fuera, me pregunté si importaría que no importara. Si nada importaba y tampoco importaba que nada importara, entonces ¿por qué no iterar? No importaba que no importara que nada importara. Y así sucesivamente.

Inhalé los vapores del peaje de Nueva Jersey y volví a considerar la situación. Si nada importaba, pero importaba que nada importara, entonces estaríamos en una situación más bien desalentadora. Si nada importaba y tampoco importaba que nada importara, entonces tendríamos la posibilidad de algo mejor —un enfoque irónico y posiblemente feliz de la vida—. Y análogamente a niveles superiores. Razonando formalmente y en un modo probablemente simple, la mejor situación sería que las cosas importaran al nivel elemental o, si no, que no importaran a ningún nivel: o la ingenuidad total de la infancia o la completa ironía del adulto. (Véase la entrada sobre *Tiempo*).

Mientras me aproximaba a la refinería Hess, mis pensamientos pasaron al tema de escribir y publicar, pero mi predisposición hacia el absurdo persistía. Dada la numerosa, y cada vez más homogénea, población lectora ¿había hoy menos «necesidad» de autores? Suponiendo que la gente lea hoy aproximadamente el mismo número de libros, revistas y diarios que en cualquier otra época, y que quieran leer siempre algo «mejor» que cualquier otra cosa que ya hayan leído (según el patrón de las listas de éxitos, por ejemplo), y que tiendan en general a leer cosas escritas por sus paisanos, parece deducirse que cuanto mayor sea un país, menor es el porcentaje de sus ciudadanos que puedan ser autores o, lo que algún día podría ser equivalente, autores de éxito.

Pensé en varios contraejemplos, el más interesante de los cuales apuntaba a la gran variedad de publicaciones (especialmente de libros y revistas no novelescos) que atienden a gustos cada vez más especializados y que proporcionan mayores oportunidades a los escritores. Si estas vagas reflexiones tenían algún sentido, la probabilidad de alcanzar el estrellato literario se reducía, mientras que aumentaban las oportunidades de ganarse la vida con un procesador de textos.

Dijeron por la radio que había una retención de una hora en el Lincoln Tunnel, así pues decidí entrar en Manhattan por el puente George Washington. Esta solución resultó ser peor al fin y al cabo, pues las víctimas de un pequeño accidente estaban en el arcén y provocaban la curiosidad habitual en los conductores que pasaban. El efecto acumulado de la reducción de velocidad de cada uno para ver que en realidad no había nada que ver me pareció una versión en miniatura de muchos problemas humanos. No había malicia, sólo un impulso común cuya ampliación tenía efectos molestos.

El tráfico se hizo fluido al cabo de unos veinte minutos sólo para volver a atascarse más aún debido a unas obras. Un tramo de carril único de un par de kilómetros antes del puente estaba innecesariamente salpicado de señales de «No Pasar». Las señales me hicieron pensar en las frases progresivas en las que cada palabra tiene una letra más que la anterior, y maté el tiempo cavilando sobre ello. Finalmente conseguí «*I Do Not Pass Since Danger Expands Anywhere Unmindful Speedsters Proliferate Unmanageably*»,^[3] *de la que me sentí desmesuradamente orgulloso.*

Cansado de esto me fijé en la frecuencia cada vez mayor de placas de médico a medida que me acercaba a Nueva York y recordé la estadística que acababa de leer de que había 428 médicos para toda Etiopía, un país de 40.000.000 de habitantes cuya esperanza media de vida es de menos de 40 años. Tratando de mantener a raya mi impaciencia, me dediqué a construir biografías de personas a partir de la vanidad de sus matrículas y sin la menor prueba llegaba cada vez a la conclusión de que acertaba. Esto me hizo pensar en el tema de un chiste de matrículas que me había contado un amigo matemático: ¿Cómo deletreas el nombre «Henry»? Respuesta: H-E-N-3-R-Y; el 3 es mudo.

Ya en el puente recordé que los cables que lo sostienen toman la forma de una curva denominada catenaria a menos que se les cuelguen pesos a intervalos iguales, en cuyo caso la forma es parabólica. Consideré la probabilidad de que el puente se derrumbara, luego la también improbable, pero mucho menos descabellada, de resultar muerto por un conductor ebrio, o por fin la de contraer un cáncer por mi repetida exposición a la autopista de peaje de Nueva Jersey, o de sufrir de hipertensión debido a la frustración de estar encerrado en un coche a solas con mis meditaciones obsesivas.

Llegué a mi cita de Nueva York con sólo cinco minutos de retraso, pero éste no es el tema de mi relato. Lo que pretendía era ilustrar un fluir matemático de conciencia. Los temas que se planteaban de un modo natural eran convenios sociales (el buen samaritano, distraerse mirando los restos de un accidente), la velocidad media (las multas por exceso de velocidad en los peajes), el nivel lógico de una frase (el tema de los «nada importa»), la probabilidad (de ser un autor publicado), los juegos con palabras (las frases «bola de nieve») y la estimación (morir por el derrumbamiento de un puente frente a otras muertes más probables).

Para la mayoría de no científicos, lo más importante de una educación científica no es la comunicación de una serie de hechos reales (aunque no pretendo con ello menospreciar el conocimiento objetivo), sino la formación de unos hábitos intelectuales científicos: ¿Cómo comprobaría eso? ¿Qué pruebas tengo? ¿Qué relación guarda con otros hechos y principios? Lo mismo vale, en mi opinión, para la educación matemática. Recordar esta fórmula o aquel teorema es menos importante para la mayoría de la gente que

su capacidad de considerar una determinada situación con una mentalidad cuantitativa, darse cuenta de las relaciones lógicas, probabilísticas y espaciales, y reflexionar matemáticamente. (Véase la entrada sobre *Cálculo y rutina*).

No pretendo decir con esto que haya que concentrarse exclusivamente en tales reflexiones, sino sólo resaltar que la matemática es mucho más que el simple cálculo, que la perspectiva que resulta de su estudio puede aclarar aspectos de nuestras vidas que están más allá de nuestras preocupaciones financieras o científicas. Y, como mínimo, nos puede proporcionar un modo alternativo de matar el tiempo mientras conducimos.

Álgebra: algunos principios básicos

El álgebra elemental siempre me gustó, en parte porque mi primera profesora había sido reclutada (quizá sería más apropiado decir forzada) para impartir la asignatura pese a no ser capaz de reconocer una ecuación de segundo grado en un ejercicio de ciclo superior. Como, no obstante, era honesta, nunca trató de ocultarlo, y al estar próxima a la jubilación, se apoyaba en sus mejores estudiantes para salir de cualquier atolladero matemático. A menudo encontraba algún pretexto para que uno de nosotros fuera a su clase al acabar la escuela, y siempre se las arreglaba para ensayar la lección del día siguiente. Por suerte (y por necesidad) insistía en unos pocos principios básicos y dejaba la mayoría de detalles para nosotros.

A pesar de saber algo de matemáticas, intentaré, en la medida de lo posible, seguir su sano ejemplo pedagógico de bosquejar una imagen general y evitar los detalles técnicos. Esto es especialmente importante en álgebra, cuya sola mención evoca en muchas personas el desdichado recuerdo de intentar determinar la edad de Enrique sabiendo que es 5 veces mayor que su hijo y que dentro de 4 años será sólo 3 veces mayor. Hay muchas razones para esta aversión, pero a veces me pregunto si una de ellas no se remonta al mismo descubrimiento del álgebra por Al-Juarizmi, a principios del siglo IX. Este Al-Juarizmi, de cuyo nombre procede la palabra «algoritmo» y de cuyo libro *Al-jabr wa'l Muqabalah* proviene la palabra «álgebra», fue uno de los mayores matemáticos de una de las épocas más impresionantes del saber árabe. Su libro trata de la resolución de varias clases sencillas de ecuaciones, pero fiel al significado que ha adquirido la palabra algoritmo, Al-Juarizmi se concentró casi exclusivamente en recetas, fórmulas, reglas y procedimientos. A mi modo de ver, este texto carece de la elegancia y el atractivo lógico de

los *Elementos* de Euclides, pero, como éste, fue la obra de referencia por excelencia en su campo durante muchísimo tiempo.

Aunque Al-Juarizmi no usara variables en sus problemas de álgebra, por la sencilla razón de que éstas no se inventarían hasta 750 años después, se ha venido a considerar el álgebra como una generalización de la aritmética en la que se usan variables en el papel de números desconocidos. (Véase la entrada sobre *Variables*). Esto permite un campo de acción muchísimo mayor pues se puede expresar, por ejemplo, la propiedad distributiva por medio de la simple igualdad $A(X + Y) = AX + AY$, mientras que en aritmética sólo se pueden citar ejemplos concretos de esta propiedad: $6(7 + 2) = (6 \times 7) + (6 \times 2)$, o $11(8 + 5) = (11 \times 8) + (11 \times 5)$. (Permítaseme intercalar aquí un conocido acertijo cuya solución se basa en la propiedad distributiva. Tómese un número entero X . Añádasele 3. Dóblese el resultado y réstesele 4. Réstese luego dos veces el número escogido y súmese 3 al resultado. Da 5. ¿Por qué?).

El título *Al-jabr wa'l Muqabalah* significa algo así como «restauración y equilibrio» y se refiere a la idea básica del álgebra de que para resolver una ecuación se han de «equilibrar» ambos miembros de la misma, y que si uno realiza una operación en un miembro de la ecuación, ha de «restaurar» la igualdad realizando también la misma operación en el otro.

Más adelante haremos un repaso de este proceso, pero como probablemente le dé lo mismo que Enrique tenga 20 años y fuera padre a los 16, considere el siguiente problema práctico sacado de la laberíntica complejidad de mi imperio financiero. Recientemente estaba tratando de decidir dónde invertir una cierta cantidad de dinero por un corto plazo de 3 meses: en un fondo que pagaba el 9% en el primer mes y luego revertía al tipo de interés de Hacienda en los dos meses restantes, o en un fondo que pagaba el 5,3% libre de impuestos. Me preguntaba cuál habría de ser el tipo de interés medio, R , de Hacienda para salir sin ganar ni perder con el fondo libre de impuestos. Esto me llevó a la ecuación $0,72[(0,09 + R + R)/3] = 0,053$, donde el 0,72 reflejaba mi nivel tributario. Ecuaciones parecidas se dan en muchos aspectos de los negocios, la ciencia y la vida cotidiana, y las sencillas técnicas empleadas para resolverlas me permitieron obtener $R = 6,54\%$, muy por debajo del tipo de interés de

Hacienda en aquel momento, con lo que invertí mi dinero en el primero de los fondos.

El álgebra trata también de los métodos de resolución de ecuaciones de segundo grado ($X^2 + 5X + 3 = 0$, por ejemplo), ecuaciones cúbicas ($X^3 + 8X^2 - 5X + 1 = 0$) y ecuaciones de grado superior ($X^N + 5X^{(N-1)} \dots - 11,2X^3 + 7 = 0$) así como ecuaciones de dos o más variables y sistemas de ecuaciones. (Véanse las entradas sobre *La fórmula de la ecuación de segundo grado* y *Programación lineal*). En todas ellas se emplean variantes y refinamientos del principio fundamental de «restauración y equilibrio» de Al-Juarizmi como la idea fundamental de que las variables representan números y, por tanto, las manipulaciones que se hagan con ellas han de obedecer a las mismas reglas que rigen a los números.

Hay otra rama de la matemática que atiende al nombre de álgebra (a veces se la llama álgebra abstracta para distinguirla del álgebra elemental), pero su lógica, sus orígenes históricos y su tono son tan distintos que reservo su discusión para la entrada sobre *Grupos*. Los algebristas cuya especialidad es el estudio de estructuras algebraicas abstractas a menudo se sienten un poco molestos cuando los legos se imaginan que su ocupación consiste principalmente en resolver ecuaciones de segundo grado.

[Deducción de la edad de Enrique sabiendo que es 5 veces mayor que su hijo y que dentro de 4 años será sólo 3 veces mayor. Empecemos suponiendo que X es la edad actual del hijo de Enrique. (Un amigo mío cuenta que empezó así una explicación y alguien en el fondo de la clase replicó sarcásticamente: «Bueno, ¿y si suponemos que no lo es?»). Entonces $5X$ ha de ser la edad actual de Enrique. Dentro de 4 años el hijo tendrá $(X + 4)$ años, mientras que Enrique tendrá $(5X + 4)$. Nos dicen que en ese momento el padre será sólo 3 veces mayor que el hijo. Expresada algebraicamente, la relación *se traduce* en la ecuación $5X + 4 = 3(X + 4)$. Ahora hemos de usar la propiedad distributiva para escribir $3(X + 4)$ como $3X + 12$. De donde $5X + 4 = 3X + 12$. Ahora pasamos a la restauración y equilibrio. Para simplificar la ecuación manteniendo la igualdad restamos 4 de ambos miembros y obtenemos $5X = 3X + 8$. Por la misma razón restamos también $3X$ de ambos miembros, obteniendo $2X = 8$. Por fin, dividimos ambos

miembros por 2 y llegamos a $X = 4$. Y como X es la edad del hijo, concluimos que $5X = 20$ es la edad de Enrique].

[La solución del otro acertijo: Si X es el número original, los resultados de las sucesivas operaciones son $(X + 3)$; $2(X + 3)$ o $(2X + 6)$; $(2X + 2)$; 2; 5].

Áreas y volúmenes

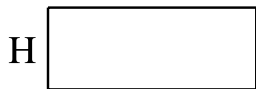
¿H abéis leído alguna vez una larguísima saga generacional y casi al final os habéis encontrado con un personaje que suelta un tópico absolutamente vacío que, a pesar de todo, encaja tan profundamente con el resto de la obra que os hace repasarla toda bajo esa nueva luz? Mi reacción ante la fórmula del área de un rectángulo es un tanto similar. Soy perfectamente consciente de su simplicidad y, al mismo tiempo, también de sus conexiones con espléndidos filones de oro matemático. Dejemos de lado el oro, sin embargo, y centrémonos en un color más pálido: el crema. Tengo ante mí un sobre de color crema; sus dimensiones son 25 centímetros por 32 centímetros y, por tanto, su área son 800 centímetros cuadrados. Como soy sumamente listo, no he contado 800 cuadrados de un centímetro de lado para encontrar esta área; me ha bastado con multiplicar la longitud del sobre por su anchura.

Es bastante simple, pero cualquier otra fórmula del área de una figura plana es una consecuencia del hecho tan sencillo de que el área de un rectángulo sea igual a la longitud de su base por la de su altura; o escrito como una ecuación, $A = BH$. A continuación siguen algunos ejemplos importantes, un par de ellos, como el anterior, tan antiguos como el propio conocimiento matemático (anteriores a los egipcios).

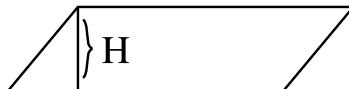
Como la base y la altura de un cuadrado son iguales, la fórmula del área del cuadrado es $A = L^2$, siendo L la longitud del lado del cuadrado. (Las fórmulas correspondientes a los perímetros del rectángulo y el cuadrado son respectivamente, $P = 2B + 2H$ y $P = 4L$).

Aunque su nombre pueda evocar un par de pesas pequeñas en paralelo, un

paralelogramo es una figura de cuatro lados —un cuadrilátero— cuyos lados opuestos son paralelos. Su área también se expresa con la fórmula $A = BH$, donde B es como antes la longitud de la base, o lado de abajo, pero ahora H es la distancia más corta entre el lado de arriba y el de abajo (o su prolongación si hiciera falta), esto es, la altura perpendicular, en contraposición con la altura inclinada.



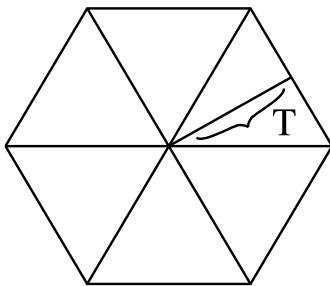
$$A = BH$$



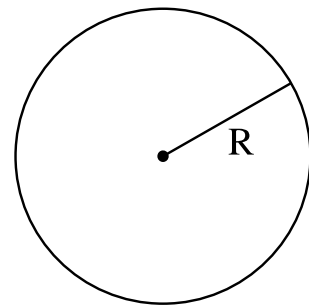
$$A = BH$$



$$A = \frac{1}{2}BH$$



Pero a medida que aumenta el número de lados del polígono inscrito, T se aproxima al radio de la circunferencia y el perímetro P se acerca al perímetro de la circunferencia C



Como el hexágono puede dividirse en triángulos, su área es igual a $\frac{1}{2}TP$, donde T es la distancia del centro a uno de los lados y P es el perímetro

Así pues, el área del círculo es igual a $\frac{1}{2}R(2\pi R)$ o, como es más corriente, πR^2 , $A = \pi R^2$

Como cualquier triángulo se puede considerar como la mitad de un rectángulo o de un paralelogramo (un corte triangular de la pieza rectangular entera), la fórmula del área del triángulo es simplemente $A = \frac{1}{2} \times B \times H$, donde, como antes, B es la base del triángulo y H la distancia más corta desde el vértice superior a la base (o su prolongación, si hiciera falta).

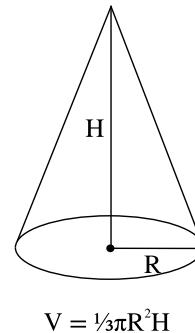
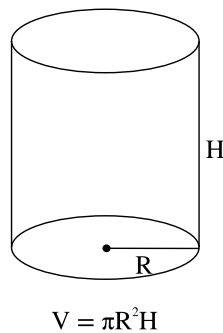
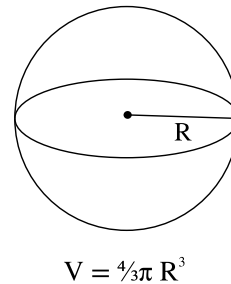
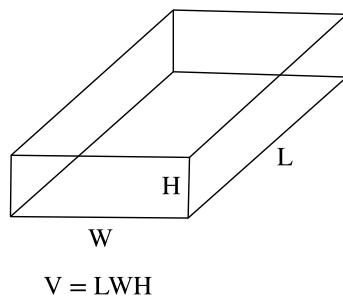
Si se empapela una habitación de planta irregular o se pintan las paredes

de una catedral gótica, se puede encontrar el área de cualquier figura plana limitada por líneas rectas, dividiendo y venciendo: se divide la figura en triángulos y rectángulos, se halla el área de cada una de estas partes y luego se suman las áreas para obtener el área de la figura global. Aplicando este método a los polígonos regulares (polígonos como los triángulos, cuadriláteros, pentágonos o dodecágonos con todos los lados y ángulos iguales), se obtiene la fórmula $A = 1/2 \times TP$, donde P es el perímetro o longitud del contorno del polígono y T es la distancia perpendicular del centro del polígono a un lado.

Además, ya que un círculo se puede considerar como la figura límite a la que converge una sucesión de polígonos regulares inscritos, también podemos usar las ideas de partir y de límite para demostrar que el área del círculo también viene dada por $A = 1/2 \times TP$, donde, como antes, T es la distancia del «punto central» o centro del círculo a un «lado» del círculo, y P es el «perímetro» del círculo. Como el perímetro del círculo es su circunferencia, 2π veces el radio R del círculo (véase la entrada sobre *Pi*), y como T es igual a R, tenemos que, cuando sustituimos estos valores en la fórmula anterior, llegamos a la expresión más conocida de $A = 1/2 \times R \times 2\pi R$ o $A = \pi R^2$ para el área de un círculo de radio R (lo que explica, entre otras cosas, por qué una pizza de 20 centímetros es casi un 80% mayor que una de 15 centímetros).

En general, para encontrar el área de una figura delimitada por líneas curvas de cualquier clase, se ha de aproximar la figura en cuestión por otra cuyo contorno esté formado por rectas. Luego se ha de calcular el área de esa figura de lados rectilíneos descomponiéndola en triángulos y rectángulos y sumando las áreas de éstos. Para tener un resultado más aproximado, se toman los segmentos rectilíneos de la aproximación lo más cortos posible, de modo que se adapten mejor al contorno curvilíneo, y luego se procede como antes por descomposición en triángulos y suma. Este «método de exhaustación» o de aproximación sucesiva de un área curva por medio de rectángulos y triángulos se remonta a Arquímedes y es la idea fundamental de la integral definida, que se define como el valor límite de estas sumas aproximadas e, incidentalmente, un concepto que se encuentra en la base de buena parte de la aplicabilidad del cálculo y del análisis matemático.

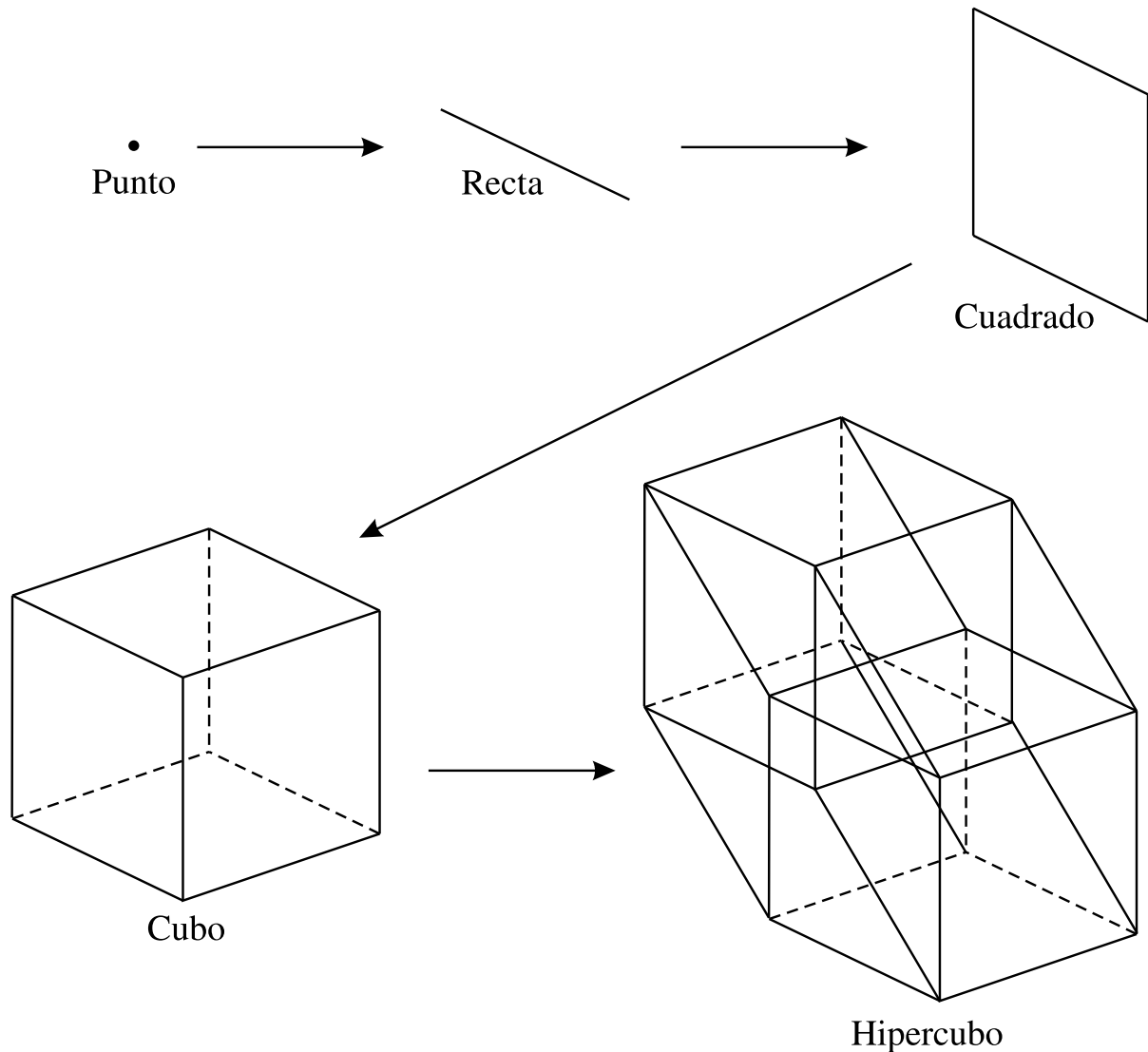
Para los volúmenes de las figuras en el espacio se puede seguir una estrategia semejante. La pieza fundamental de esas figuras es la caja rectangular, cuyo volumen V se halla multiplicando la longitud por la anchura por la altura, $V = LAH$; $V = L^3$ para el cubo de lado L . El volumen de un paralelepípedo, figura de seis caras cuyas caras opuestas son paralelogramos iguales, viene dada por la misma fórmula pero, como antes, la anchura y la altura se han de interpretar como distancias perpendiculares entre caras opuestas. Otras fórmulas útiles son las de un bote cilíndrico (expresión tan redundante como «nueva innovación»): $V = \pi R^2 H$, siendo R el radio del bote y H su altura; de un cono, $1/3 \times \pi R^2 H$, donde R es el radio de la base del cono y H su altura perpendicular; y de una esfera, $4/3 \times \pi R^3$, donde R es el radio de la esfera.



Fórmulas de los volúmenes elementales

Como en el caso de las áreas, el volumen de una figura más compleja se puede determinar calculando el límite de sucesivas aproximaciones a la figura por medio de cajas rectangulares; esto es, usando de nuevo el método de

exhaustación de Arquímedes, formalizado en el cálculo integral. No podemos pasar sin citar otras generalizaciones naturales de las fórmulas del área y el volumen a hiperespacios de más dimensiones (en los que los hipercubos serían las piezas elementales), así como otros temas más teóricos relativos a la naturaleza y propiedades de las áreas y los volúmenes (de superficies, sólidos y conjuntos arbitrarios).



Representación bidimensional de un hipercubo

Entre los conceptos de área, volumen y física elemental se da una curiosa interacción. Obsérvese que la fuerza de sustentación que necesita una

persona, un animal o una estructura cualquiera para mantenerse en pie, es proporcional al área de su sección transversal, mientras que su peso es proporcional al volumen. Por ejemplo, si se cuadruplica la altura de una estructura conservando sus proporciones y el material de que se compone, se tiene que el peso se multiplica por 64 (4^3), mientras que su capacidad para soportar peso sólo se ha multiplicado por 16 (4^2). Ésta es la razón por la que cualquier monstruo de 8 metros de altura que se pasee por el Himalaya o se tueste al sol en alguna playa del Triángulo de las Bermudas no puede tener las mismas proporciones que nosotros. Esta relación también condiciona las alturas, proporciones y materiales de los árboles, edificios y puentes. Consideraciones semejantes permiten explicar otras características estructurales (como el área de la superficie del pulmón y del intestino) de las plantas, animales y objetos inanimados. (Véase la entrada sobre *Fractales*).

Repito lo que dije al principio de este artículo. Aunque la fórmula $A = BH$ sea en cierto sentido trivial (la ilustración arquetípica de la multiplicación), toda la matemática está llena de variantes, refinamientos y aplicaciones de esta idea.

Finalmente, el saber estas fórmulas de las áreas y los volúmenes no siempre garantiza un sentido intuitivo de la extensión y la voluminosidad. Baste con un ejemplo: el Gran Cañón del Colorado tiene 369 kilómetros de largo, su anchura oscila entre 7 y 31 kilómetros, y llega a tener hasta 1,7 kilómetros de profundidad. Si, siendo un poco conservadores, le asignamos 10 kilómetros de anchura media y 0,5 kilómetros de profundidad media, su volumen es de 1.845 kilómetros cúbicos, que multiplicado por 1.000^3 da $1,845 \times 10^{12}$ metros cúbicos. Si dividimos esta cifra por cinco mil millones, el total de la población humana de la Tierra, obtenemos aproximadamente 369 metros cúbicos de espacio del Gran Cañón por cada ser humano. Sacando la raíz cúbica de esta cifra, aproximadamente 7 metros, llegamos a que en la Gran Colmena cabrían cinco mil millones de cubículos de 7 metros de lado.

Cálculo

El cálculo infinitesimal, descubierto independientemente a finales del siglo XVII por Isaac Newton y Gottfried Wilhelm von Leibniz, es la rama de la matemática que trata de los conceptos fundamentales de límite y variación. Al igual que la geometría axiomática de los antiguos griegos, casi desde sus comienzos tuvo un profundo efecto en el pensamiento científico, matemático y del público en general. Ello se debe parcialmente a la potencia, elegancia y versatilidad de sus técnicas y también a su asociación con la física newtoniana y a la metáfora del universo como mecanismo gigantesco gobernado por el cálculo y unas ecuaciones diferenciales eternas. En las últimas décadas los avances de la física moderna han quitado bastante fuerza a esta imagen, a la vez que los progresos de los ordenadores han cuestionado la posición, otrora indiscutida, del cálculo en los planes de estudio. A pesar de todo, el cálculo sigue siendo una de las ramas más importantes de la matemática para el científico y el ingeniero, y cada vez más para el economista y el hombre de negocios.

Aunque sea un poco simplista, es útil dividir la materia en dos partes: el cálculo diferencial, que trata de las tasas de variación, y el cálculo integral, que se ocupa de sumar cantidades que varían. Empezando por el cálculo diferencial, como hacen la mayoría de tratados, supongamos que después de comer salimos de Filadelfia en dirección a Nueva York por la autopista de Nueva Jersey. El coche lleva reloj y cuentakilómetros, pero no velocímetro, aunque notamos que la velocidad cambia debido a la intensidad del tráfico, la música o el estado de ánimo del conductor. Una pregunta que se nos plantea de manera natural es: ¿cómo podríamos determinar la velocidad instantánea en un momento dado, la una, por ejemplo? Supongamos que nos interesa más

una definición teórica que una manera práctica de realizarlo.

Una primera respuesta aproximada podría consistir en tomar la velocidad media entre la 1:00 y la 1:05. Si recordamos que la velocidad media es igual a la distancia recorrida dividida por el tiempo empleado, podríamos usar el cuentakilómetros para determinar la distancia recorrida en estos cinco minutos y luego dividirla por $1/12$ (un doceavo) de hora (cinco minutos). Tendríamos una mejor aproximación si determináramos nuestra velocidad media entre la 1:00 y la 1:01, buscando la distancia recorrida durante este minuto y dividiéndola luego por $1/60$ de hora. Este segundo resultado se aproximaría más a la velocidad instantánea a la 1:00, pues dejaría menos tiempo para posibles cambios de velocidad. Obtendríamos una aproximación mejor aún si encontráramos la velocidad media entre la 1:00:00 y la 1:00:10. Como antes, determinaríamos la distancia recorrida durante este intervalo de diez segundos y la dividiríamos por $1/360$ de hora.

No se trata de un método demasiado eficaz, pero sirve para llevamos a la definición teórica de velocidad instantánea. La velocidad instantánea en un instante dado es, por definición, el límite de las velocidades medias sobre intervalos de tiempo cada vez más cortos que contengan el instante en cuestión. Empleado aquí, «límite» es un término delicado (véase la entrada sobre *Límites*), pero me parece que en este caso su aplicación es bastante intuitiva. Además, y esto es importante, si la distancia que hemos recorrido por la autopista de Nueva Jersey viene dada por una fórmula que depende sólo del tiempo que llevamos viajando, el cálculo nos proporciona técnicas que nos permiten determinar la velocidad instantánea a partir de dicha fórmula. Si representáramos gráficamente esta fórmula que relaciona la distancia recorrida (sobre el eje Y) con el tiempo empleado (sobre el eje X), la velocidad en cualquier instante correspondería a lo empinado de la gráfica en el punto dado, es decir, a su pendiente en ese punto.

La definición y las técnicas son muy generales y son las que aparecen de un modo natural siempre que uno se plantea la pregunta genérica: ¿a qué velocidad está cambiando esto? Como en el caso anterior, a menudo nos interesa saber cómo cambia una cierta cantidad con el transcurso del tiempo. ¿A qué velocidad conducíamos a la una? ¿A qué velocidad se estará extendiendo la mancha de petróleo al cabo de tres días? ¿A qué velocidad se

alargaba la sombra hace una hora? Aunque a menudo nos interesan también otras tasas de cambio más generales. ¿Cuánto aumentarán nuestros beneficios con respecto al número de artículos fabricados si producimos 12.000 diarios? ¿Cuánto aumentará la temperatura de un gas contenido en un recipiente con respecto al volumen, si éste es de 5 litros? ¿Cuánto aumentarán las ganancias con respecto al capital invertido si éste es de 800 millones de dólares (suponiendo que los demás factores se mantienen constantes)? Siempre que se conozca la relación entre las cantidades implicadas, las técnicas del cálculo diferencial se pueden usar para determinar la tasa de variación —que se conoce como la «derivada»— de una cantidad con respecto a la otra. [Si la relación entre las cantidades x e y viene dada por la fórmula $y = f(x)$ (véase la entrada sobre *Funciones*), entonces la derivada se indica con una fórmula que se suele simbolizar como $f'(x)$, y que en la notación de Leibniz se escribe dy/dx . La fórmula de la derivada nos dice a qué ritmo varía y con respecto a x en cualquier punto x .

Como suele ocurrir en matemáticas, saber las fórmulas, en este caso las de las derivadas obtenidas por estas técnicas, no tiene por sí mismo mucho valor. Todos los estudiantes de cálculo «saben» que la derivada de $Y = X^N$ es NX^{N-1} . Para demostrar la superficialidad de este conocimiento, cuando mis hijos iban al parvulario les enseñé a contestar NX^{N-1} siempre que les preguntaba cuál era la derivada de X^N . Ellos también «sabían» cálculo.

Muchos tipos de problemas resultan fáciles una vez se ha comprendido el concepto de derivada. Como los beneficios, por ejemplo, acostumbran a aumentar y después a disminuir en función del número de artículos producidos, sabemos que la tasa de variación del beneficio con respecto a los artículos primero es positiva (aumento del beneficio con el número de artículos producidos) y luego negativa (disminución de los beneficios si seguimos fabricando más). Si conocemos la relación precisa entre beneficios y artículos podemos determinar cuántos productos hemos de fabricar para maximizar nuestros beneficios, encontrando cuándo se anula el valor de la tasa de variación. Podemos emplear la misma técnica para optimizar recursos escasos.

Para hacerse una ligera idea de lo que es el cálculo integral, supongamos

que estamos otra vez en la autopista de Nueva Jersey (el camino real al cálculo), pero que esta vez el coche está equipado de un reloj (pongamos que marca las 2:00) y un velocímetro, pero no tiene cuentakilómetros. La monotonía de conducir nos ha hecho caer en un talante reflexivo y nos preguntamos cómo podríamos saber la distancia recorrida durante la próxima hora en función de la velocidad. Si mantenemos una velocidad constante de 80 kilómetros por hora, el problema es trivial: habremos recorrido exactamente 80 kilómetros.

Sin embargo, como la velocidad de nuestro coche cambia considerablemente, podríamos intentar dar una respuesta aproximada del modo siguiente: consultemos el velocímetro a las 2:02:30 y supongamos que la velocidad (pongamos 96 km/h) se mantiene constante en el intervalo entre las 2:00 y las 2:05. Como la distancia recorrida en cualquier intervalo de tiempo es igual al producto de la velocidad por la duración de dicho intervalo, multiplicamos 96 km/h por $1/12$ de hora, con lo que obtenemos un valor aproximado de la distancia recorrida entre las 2:00 y las 2:05 (unos 8 kilómetros). Consultemos a continuación el velocímetro a las 2:07:30 y supongamos que nuestra velocidad (pongamos que ahora es de 80 km/h) permanece aproximadamente constante durante el intervalo que va de las 2:05 a las 2:10, multiplicamos 80 km/h por $1/12$ de hora y tenemos un valor aproximado de la distancia recorrida entre las 2:05 y las 2:10. Quizás a las 2:12:30 hayamos encontrado un tráfico muy denso y hayamos tenido que reducir a 56 km/h. Multiplicamos esta velocidad por $1/12$ de hora y tendremos una estimación de lo que hemos recorrido entre las 2:10 y las 2:15. Siguiendo así, sumemos todas estas distancias y obtendremos la distancia total recorrida durante la hora.

La variación de la velocidad de nuestro coche será menor en un minuto que en cinco. Por tanto, si queremos una aproximación mejor, emplearemos intervalos de un minuto en vez de cinco y, como antes, sumaremos todos los pequeños trozos de distancia. O también podríamos sumar todas las distancias recorridas en intervalos de diez segundos sucesivos, con lo que obtendríamos un resultado más aproximado para la distancia recorrida durante esa hora. La distancia recorrida exacta se define como el límite de este procedimiento, y dicho límite se conoce con el nombre de «integral

definida» de la velocidad. El resultado de la suma dependerá, naturalmente, de la velocidad y del modo exacto como ésta varíe a lo largo de la hora.

Como en el caso de las tasas de variación, el procedimiento es completamente general y se presenta cuando uno se pregunta acerca de una cantidad variable: ¿A cuánto asciende en total? Por ejemplo, una aproximación de la fuerza total que ejerce un embalse sobre la presa que lo contiene la podemos obtener sumando la fuerza contra el estrato inferior de un metro de altura a la fuerza contra el estrato de un metro de altura inmediatamente superior, luego a la fuerza sobre el estrato siguiente y así sucesivamente hasta llegar a la parte superior de la presa. Hemos de hacerlo así porque la presión del agua, y por tanto la fuerza que ejerce, aumenta con la profundidad. Se obtiene una mejor aproximación dividiendo la presa en estratos de un centímetro de altura y sumando las fuerzas que actúan sobre cada uno de ellos, y la fuerza exacta se obtiene encontrando el límite de este procedimiento —la integral definida—. Análogamente, si tratamos de encontrar los ingresos totales resultantes de vender productos cuyo precio disminuye continuamente con la cantidad fabricada, vamos a parar también al concepto de integral definida. [La integral de una cantidad $y = f(x)$ se suele indicar por $\int f(x) dx$, donde el signo $\int$ es una S estilizada que significa «suma»].

La utilidad de la integral indefinida procede en gran parte del llamado teorema fundamental del cálculo, según el cual esta operación y la otra operación fundamental que hemos presentado, la de hallar la tasa de variación o derivada de una cantidad respecto a otra, son de hecho operaciones inversas, es decir, cada una deshace los efectos de la otra. El teorema y las técnicas que se desprenden de las dos definiciones nos proporcionan los útiles necesarios para comprender las cantidades que varían continuamente. Las ecuaciones diferenciales (ecuaciones en las que aparecen derivadas —véase la entrada correspondiente—) son un ejemplo particularmente valioso de la aplicación de estos útiles.

Estas ideas estimularon en gran manera el desarrollo del análisis matemático; el cálculo y las ecuaciones diferenciales se convirtieron en el lenguaje de la física y el mundo cambió para siempre. Recuérdelo la próxima vez que conduzca por una autopista con un velocímetro o un

cuentalómetros estropeado.

Cálculo y rutina

Ambition, Distraction, Uglification y Derision [«Ambición, distracción, afeamiento e irrisión»] son los nombres que daba Lewis Carroll a las cuatro operaciones básicas de la aritmética: adición, sustracción, multiplicación y división. Ésta es aún la idea que conserva mucha gente, yo entre ellos, del cálculo aritmético de la edad escolar (exceptuando la «ambición», que nunca pareció pertenecer a la lista y que quizá debiéramos sustituir por «adicción»). Las razones de esta aversión son, a mi entender, que el cálculo es aburrido, cansado y agobiante. Peor aún, a menudo pone color (o más bien lo quita) a la imagen que la gente tiene de la verdadera matemática.

Imaginemos por un momento que el 90% de cada curso de lengua, desde la enseñanza primaria hasta la universidad, se dedicara a estudiar la gramática y a analizar frases. ¿Tendrían los licenciados alguna apreciación por la literatura? O consideremos un conservatorio en el que se dedicara el 90% de los esfuerzos a practicar escalas. ¿Se formaría en los estudiantes una apreciación y una comprensión de la música? La respuesta es no, desde luego, y, sin embargo, con una cierta licencia para la hipérbole, esto es lo que ocurre frecuentemente en nuestras clases de matemáticas. Las matemáticas se identifican con un recitado rutinario de hechos y una ciega aplicación de métodos. Décadas después este modo robótico de comportarse vuelve siempre que se plantea un problema matemático. Casi todo el mundo siente que si no les dan la respuesta, o por lo menos una receta para encontrarla, nunca lo sabrán resolver. La idea de pensar sobre un problema o de discutirlo con alguien más les resulta completamente nueva. ¿Pensar sobre un problema de mates? ¿Discutirlo? (Véanse también las entradas sobre *Sustituibilidad* y

Humor).

En mi opinión la atención que la escuela da al cálculo es excesiva y obsesiva. Naturalmente, no hay nada malo en saber las tablas de sumar y multiplicar, así como en conocer los algoritmos elementales para tratar con fracciones, porcentajes, etc. De hecho, conocer estas técnicas es esencial aún hoy en día, cuando con una calculadora de mil pesetas (una parte del material escolar de cualquier niño) se pueden hacer todos los cálculos que hagan falta a la mayoría. Es precisamente después de *un poco* de empleo rutinario cuando estas técnicas habrían de tomarse como útiles para profundizar en la comprensión de los problemas, y no como un sustituto de esta comprensión.

Una consecuencia inapreciada de esto, que aparentemente es un tópico, es que la matemática habría de entenderse como algo unido sin solución de continuidad con el lenguaje y la lógica (véanse las entradas sobre *Variables*, *Los cuantificadores* y *Al estilo matemático*) y no como un conjunto aislado de ejercicios isométricos mentales. En la escuela primaria, por ejemplo, debería haber lecciones dedicadas a decidir cuál es la operación aritmética, o la sucesión de operaciones, indicada para resolver un problema dado; a estimar magnitudes muy grandes o muy pequeñas; a relatos detectivescos con tintes matemáticos; a patrones numéricos y acertijos mecánicos (verbigracia, el cubo de Rubik); a juegos de mesa (como el Monopoly) en los que entra el azar; a aspectos matemáticos de las noticias del periódico y de los acontecimientos deportivos (medias de enceste y rebotes) y a un montón de otros temas que puedan tener relación con la vida cotidiana de un niño.

Si esta conexión entre las matemáticas y las formas de pensamiento y lenguaje corrientes se establece a una edad temprana, las tablas, fórmulas y algoritmos que vendrán después están justificados: no son más que un medio abreviado de encontrar la solución. Los llamados problemas de letra (no una clase ontológica natural, sino uno cualquiera de la infinidad de problemas matemáticos que se expresan en palabras) no serían una aberración de la clase de matemáticas, sino su núcleo primario. El absurdo lamento «puedo con las mates pero no con los problemas de letra» se oiría con menos frecuencia en los institutos y universidades. Dondequiera que lo oigo ahora me pregunto cuál cree esa persona que es el objeto de las «mates». ¿Acaso su razón fundamental son páginas de polinomios que factorizar o páginas de funciones

que derivar?

La insistencia constante en el cálculo en la escuela temprana conlleva la tiranía de la respuesta correcta, otro obstáculo para el aprendizaje de las matemáticas y otro aspecto del todavía demasiado común modelo convento-cuartelario de su enseñanza. Ésta es la Verdad; ahora resuelvan estos 400 problemas idénticos. En la mayoría de los otros campos hay una clara distinción entre respuestas incorrectas, pero demasiada gente cree que si en matemáticas una respuesta no está bien, está mal, y punto. Cuando lo cierto es precisamente lo contrario. Si dos personas tuvieran que sumar $2/5$ y $3/11$ y una diera como respuesta $5/16$ y la otra da $39/55$, estaría bastante claro que la primera no sabe mucho de quebrados, mientras que la segunda sólo ha sido poco cuidadosa. (En realidad, se podría defender incluso la primera «suma»: $2/5 + 3/11 = 5/16$. Quizás el $2/5$ significa que un jugador de baloncesto ha metido 2 canastas de 5 intentos en la primera, parte y el $3/11$, que ha encestado 3 de los 11 lanzamientos intentados en la segunda parte. En todo el partido habría encestado 5 tiros de 16, con lo que la «suma» anterior estaría justificada).

En álgebra, aritmética y probabilidad elemental normalmente hay varias maneras de resolver un mismo problema y, en problemas más difíciles y no tan bien definidos, hay más. (Yo suelo calificar con puntuación máxima las respuestas erróneas si las «matemáticas» están mal pero la concepción es correcta, y doy una puntuación parcial si el enfoque del problema es razonable). La creencia común de que todas las respuestas erróneas son equivalentes, o incluso de que todas las respuestas correctas son equivalentes, rebaja la necesidad de pensar críticamente, cosa que explica la frecuencia con que esto se da entre los estudiantes y, aunque sea triste decirlo, también entre muchos profesores.

Me da la impresión de que los lamentos acerca de la poca capacidad de nuestros chicos para hacer cálculos simples son parecidos a los debates que quizá se desencadenaron en la Italia del siglo XV acerca de las dificultades de los estudiantes con los algoritmos de la división en numeración romana. Gradualmente se vería que era difícil adquirir destreza en esta técnica concreta y que, debido al nuevo *software* de la numeración arábica, era menos útil de lo que había sido hasta el momento. De un modo atenuado ésta

es la situación actual. La habilidad para calcular a mano es menos útil que antes, y ésta es otra razón por la que habríamos de abandonar nuestra insistencia fundamentalista en la capacidad de cálculo.

[Antes de empezar el doctorado, hice una corta estancia con el Peace Corps y enseñé matemáticas en una escuela secundaria en un distrito muy pobre de Kenia occidental. A pesar de la falta de materiales, el nivel matemático de los estudiantes era considerablemente mejor que el de los estudiantes de las escuelas más ricas de Nairobi, donde disponían de un montón de libros pasados de moda, llenos de ejercicios aburridos página tras página. El poco personal de mi escuela trabajaba sobre problemas prácticos (por lo menos en principio) y en la comprensión de los conceptos. Aunque éste no sea ciertamente un estudio concluyente, la experiencia aumentó mi descontento con la pedagogía tradicional].

La matemática no sólo es cálculo, igual que escribir no es sólo mecanografía. Casi todo el mundo acaba a la larga por aprender a calcular, pero según los informes relativos a nuestra enseñanza en matemáticas, no se fomentan en nuestros chicos otras capacidades de niveles superiores. Muchos estudiantes de secundaria no saben interpretar gráficas ni entienden conceptos estadísticos, son incapaces de hacer modelos matemáticos de una situación, raramente estiman o comparan magnitudes, nunca demuestran teoremas y, lo que es más lamentable, apenas son capaces de mostrar una actitud crítica y escéptica con respecto a los datos y las conclusiones, ya sean numéricos, espaciales o cuantitativos. Los costes públicos y privados de este anumerismo y de esta incapacidad matemática general son incalculables.

Me gustaría poder decir que este énfasis en el cálculo se acaba cuando los estudiantes llegan a la universidad. Pero ¡ay!, incluso en los cursos de análisis, álgebra lineal y ecuaciones diferenciales (que se siguen en muchas carreras) se encuentra la misma tendencia atontadora a poner problemas de cálculo rutinario. Los estudiantes están acostumbrados a este enfoque; los libros de texto tienen que atraer a una gran clientela, cosa que los hace blandos, y enseñar así supone para los profesores una menor exigencia, tanto en sentido social como intelectual o de tiempo. (Esto último es importante, pues el reconocimiento que reciben los profesores universitarios de matemáticas por ser buenos educadores es, en el mejor de los casos, mínimo;

y, en cambio, son recompensados si publican, independientemente de que sus publicaciones añadan sólo el más sutil de los matices a los detalles más oscuros de sus estrechas especialidades). Este énfasis en el cálculo es de lo más desafortunado si atendemos al nuevo *software* que nos libera de todo el penoso trabajo de tareas como evaluar integrales definidas e invertir matrices. (No incluyo en esta acusación los cursos superiores de la licenciatura ni los cursos de doctorado, aunque el nivel de instrucción generalmente superior que se da en ellos es poco cómodo para la mayoría).

Es de mala educación protestar repetidamente contra una enseñanza repetitiva, de modo que resumiré. La matemática es pensar —sobre números y probabilidades, acerca de relaciones y lógica, o sobre gráficas y variaciones —, pero, al fin y al cabo, pensar. Este mensaje ni tan siquiera llega a rozar a muchos de nuestros mejores estudiantes, que siguen viéndola como «Ambición, Distracción, Afeamiento e Irrisión».

Cintas de Möbius y orientabilidad

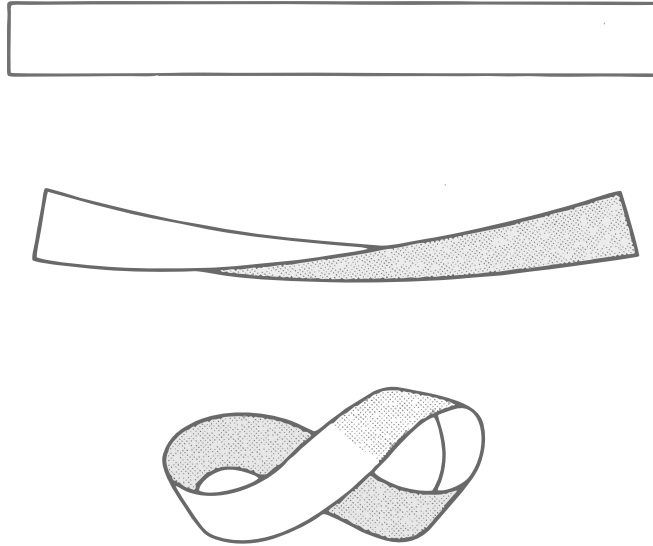
Tome una lata de atún y quítele la etiqueta con cuidado. Es una larga cinta rectangular impresa por un lado y blanca por el otro. Si damos medio giro a esta banda de papel y pegamos sus dos extremos, procurando que la cara blanca encaje perfectamente con la cara exterior impresa, obtenemos una cinta de Möbius, famosa por sus extrañas propiedades topológicas.

La principal de estas propiedades es que la cinta de Möbius tiene una sola cara. Hay un cambio continuo de blanco a impreso y otra vez a blanco. Dicho de otro modo, puedo afirmar que ni usted ni nadie podría ganar los cien millones de pesetas que alguien le prometiera por pintar una cara de la cinta de Möbius de rojo y la otra de azul. Si se empieza en rojo en cualquier punto de la cinta y se pinta sin parar, se llega irremediabilmente al mismo punto de partida habiéndola pintado toda de rojo.

Para entender otra extraña propiedad de esta figura, imaginemos una línea que la recorra por el medio. Si cortamos siguiendo dicha línea parece como si la cinta de Möbius se hubiera de romper en dos partes separadas. Pues no. El resultado no es otro que una cinta de Möbius más larga. Si en vez de ello cortamos la cinta por una línea paralela al borde pero que diste de él un tercio de su anchura en lugar de la mitad, el resultado son dos cintas enlazadas, una de ellas de Möbius.

La cinta de Möbius de una sola cara es una de las más conocidas de un gran conjunto de aberraciones topológicas. (Véase la entrada sobre *Topología*). Aunque no tenga aplicaciones importantes ni se manifieste en la naturaleza (al menos por el momento), su sorprendente simplicidad resulta atractiva. Simple como es, resulta notable que esta curiosidad no se

descubriera antes, pero el honor de su alumbramiento pertenece al astrónomo alemán del siglo XIX, A. F. Möbius.

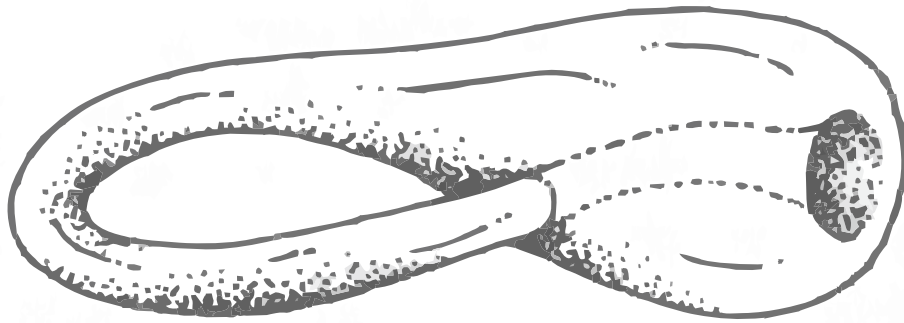


La cinta de Möbius

Möbius descubrió también que no hay ninguna manera consistente de asignar una orientación a la cinta. Para verlo mejor, imaginemos un objeto bidimensional en forma de mano y hagámoslo deslizar alrededor de una cinta de Möbius. Recordemos que la cinta de Möbius ideal no tiene grosor y, por tanto, la mano será visible desde ambas «caras» de la cinta. Observaremos que, al regresar al punto de partida, la orientación de la mano está invertida. Una mano izquierda se convierte en derecha y viceversa. Los físicos han especulado con la posibilidad de que el universo fuera «no orientable» como una cinta de Möbius (cósmicamente disléxico, si lo prefieren), de modo que, después de hacer un largo viaje cósmico, un astronauta pudiera regresar a la Tierra con el corazón en el lado derecho.

El concepto de orientación depende del número de dimensiones. Si uno recorta dos pedazos de cartón en forma de mano, una derecha y otra izquierda, y los hace deslizar sobre el suelo, no hay manera de hacerlos coincidir. Pero si levanta una de las «manos» a la tercera dimensión, basta con girarla para hacerla coincidir con la otra. La cinta de Möbius también tiene esta propiedad, pero de una manera más retorcida que refleja su peculiar

forma de sumergirse en el espacio tridimensional.



Una botella de Klein no tiene interior ni exterior. Sólo se puede realizar en un espacio de cuatro dimensiones, donde no interacciona consigo misma, como parece hacerlo en esta representación bidimensional

Un análogo tridimensional de la cinta de Möbius es la botella de Klein, que no tiene interior ni exterior. Si se corta por la mitad se obtienen dos bandas de Möbius, cada una de las cuales es imagen especular de la otra. Para hacerse una idea visual de la botella de Klein, que sólo se puede realizar en un espacio de cuatro dimensiones, hay que recurrir a los trucos normales: mirar secciones de la figura de dimensión menor (bidimensionales o tridimensionales) y examinar sus proyecciones o sombras. Los mismo trucos valen también para figuras de más dimensiones, pero al cabo de un rato uno abandona las visualizaciones y recurre a trabajar con estas figuras de una manera puramente formal, tratando las dimensiones como simples archivos matemáticos. En este sentido, un punto en un espacio de cinco dimensiones, por ejemplo, es una sucesión ordenada de cinco números, y una «hipersuperficie» es una colección de esos puntos. La no orientabilidad de tal superficie se convierte en ciertas relaciones algebraicas entre sus puntos.

Clasificar y recuperar

Parece como si clasificar y recuperar fueran técnicas menores de contabilidad que no precisaran de demasiada formación matemática. Aunque se encuentre un placer tonto en ordenar a mano alfabéticamente una lista (algo tan simple como tricotar, supongo) o en encontrar una referencia en un gran cajón de fichas, pocas personas se han detenido a pensar en los aspectos teóricos de estas actividades. Sin embargo, hallar la mejor manera de realizarlas es un problema matemático interesante que tiene muchísima importancia práctica.

Suponga que le han pedido que ordene un gran montón de papeletas. Un método podría consistir en comparar sucesivamente cada papeleta con las que ya están ordenadas, colocarla en el lugar que le corresponde y luego hacer lo mismo con la siguiente papeleta del montón. O también podría dividir el montón en muchas pilas más pequeñas que ordenaría por el método que fuera. Luego reuniría estas pilas a pares y combinaría sus ordenaciones comparando los primeros elementos, los segundos elementos, etc. Con las pilas mayores así obtenidas volvería a hacer lo mismo: aparearlas y combinar sus ordenaciones. De este modo iría disminuyendo el número de pilas ordenadas y aumentando su tamaño, hasta que al final acabaría por tener una sola pila ordenada, con lo que habría terminado su tarea.

El método elegido no tiene demasiada importancia si sólo hay unas docenas de entradas, pero si son miles o millones la diferencia entre los tiempos necesarios para uno u otro método puede ser enorme. (Estoy suponiendo que el clasificador, ya sea una persona o un ordenador, puede hacer dos cosas: comparar dos números y mover un número de un lugar a otro). El primer método, que se llama algoritmo de clasificación por

inserción, necesita, en el peor de los casos, aproximadamente N^2 pasos (o unidades de tiempo) siendo N el número de entradas a ordenar, mientras que el segundo, que se llama algoritmo de clasificación por combinación, sólo necesita unos $N \times \ln(N)$ (donde $\ln()$ representa la función logaritmo natural, véase la entrada sobre *E*) pasos para el mismo número N de entradas. Si N es 100, N^2 es 10.000, mientras que $N \times \ln(N)$ es sólo 460 —y la diferencia ya es sustancial.

Los algoritmos de recuperación diseñados para extraer y sacar fragmentos de información de una larga lista y después relacionarlos de diversas maneras consumen frecuentemente más tiempo que cualquiera de los dos algoritmos de clasificación anteriores. (Esto vale especialmente si los artículos son muy parecidos: es mucho más fácil encontrar una aguja en un pajar que en un montón de agujas). Algunos de estos algoritmos precisan de 2^N pasos para N entradas y es este hecho en particular el que nos convence de la importancia práctica de estas ideas. Si tomamos otra vez $N = 100$, 2^N es aproximadamente $1,3 \times 10^{30}$, un número tan enorme de pasos que hace que el algoritmo sea prácticamente inútil (y también inútil en la práctica). No es descabellado pensar que el fracaso de las economías centralizadas se pueda deber tanto a condicionantes de la teoría de la información como a condicionantes políticos, al encontrarse los comisarios con una dificultad creciente en la coordinación centralizada de unos datos exponencialmente crecientes acerca de la oferta, las partes y la logística. (Véase también la entrada sobre *La complejidad*).

El problema es universal. Ahora que una impresora láser puede convertir un ordenador personal en un centro de publicaciones o en una fundición de letra impresa, nuestra capacidad de clasificar y recuperar información ha caído todavía más por debajo de nuestra capacidad de producirla. A medida que crece rápidamente la cantidad de boletines financieros y artículos de investigación, de noticias y periódicos, de bases de datos y el volumen del correo electrónico, o de libros de texto o de cualquier otra clase, el número de sus interdependencias crece exponencialmente. Necesitamos nuevas maneras de interrelacionarlos, de encontrar referencias cruzadas y de determinar prioridades si no queremos anegarnos en un mar de información.

Frecuentemente tenemos más información de la que somos capaces de manejar. El informático Jesse Shera da tristemente en el clavo cuando dice, parafraseando a Coleridge: «Datos, datos por todas partes, pero ni una sola idea para pensar». Cada vez es mayor el número de personas que se basan solamente en resúmenes, reseñas, sumarios y estadísticas, pero carecen de los útiles conceptuales necesarios para llenarlos de contenido. El algoritmo de clasificación más importante que existe es una buena formación y una amplia cultura general.

Coincidencias

Las coincidencias nos fascinan. Parece como si nos obligaran a buscarles un significado. Sin embargo, más a menudo de lo que alguna gente piensa, son completamente esperables y no precisan una explicación especial. Seguramente no se puede extraer ninguna conclusión cósmica del hecho de que hace poco y por pura casualidad me encontrara a alguien en Seattle cuyo padre había jugado en el mismo equipo de béisbol del instituto de Chicago que el mío, y cuya hija tiene la misma edad y se llama igual que la mía. Por improbable que fuera este suceso *particular* (como lo son siempre los sucesos particulares), es muy probable que *algún* suceso de esta clase tan vagamente definida se produzca de vez en cuando.

Concretando más, puede demostrarse, por ejemplo, que si dos extraños se sientan juntos en un avión, más del 99% de las veces estarán unidos de alguna manera por dos o menos intermediarios. (La relación con el compañero de curso de mi padre era más sorprendente. Sólo había un intermediario, mi padre, y contenía otros elementos). Quizá, por ejemplo, el primo de uno de los pasajeros conozca al dentista del otro. La mayoría de las veces la gente no descubre estas relaciones porque en una conversación casual nadie suele hacer un repaso de sus aproximadamente 1500 conocidos ni de los conocidos de sus conocidos. (Imagino que al popularizarse cada vez más los ordenadores de sobremesa podrían comparar sus respectivas bases de datos personales y también los de las personas conocidas. Quizás intercambiar bases de datos podría convertirse pronto en algo tan corriente como dejar la tarjeta de presentación. Tejiendo una red electrónica. Infernal).

Sin embargo, hay una tendencia a buscar conocidos comunes. Tales conexiones se descubren pues con una frecuencia suficiente, de modo que los

chillidos de sorpresa que siguen a esos descubrimientos son injustificados. Igual de poco convincente es el sueño «profético» que tradicionalmente sale a la luz después de que se haya producido algún desastre natural. Si tenemos en cuenta que Estados Unidos tiene quinientos millones de horas de sueño cada noche —2 horas por noche por 250 millones de personas— es perfectamente esperable.

O consideremos también el famoso problema del cumpleaños en teoría de la probabilidad. Habría que reunir 367 personas (una más que los días de un año bisiesto) para estar seguros de que al menos dos de ellas celebran el cumpleaños en el mismo día. Pero si se quiere tener sólo una probabilidad del 50% de que esto ocurra basta con reunir 23 personas. En otras palabras, si imaginamos una escuela con miles de clases, cada una de las cuales tiene 23 alumnos, entonces aproximadamente la mitad de las clases tiene dos estudiantes que nacieron el mismo día. No hay que perder ni un minuto en tratar de explicar el significado de éstas u otras coincidencias. Simplemente ocurren.

Un ejemplo un poco distinto es el del editor de un boletín bursátil que manda 64.000 cartas en las que ensalza las posibilidades de su base de datos, sus contactos y sus sofisticados modelos econométricos. En 32.000 de estas cartas predice una alza de determinado índice bursátil para la semana siguiente, y en las 32.000 restantes predice una baja del mismo índice. Ocurra lo que ocurra, manda una segunda carta, pero sólo a los 32.000 que recibieron una «predicción» correcta. En 16.000 de ellas predice una alza para la semana siguiente y en 16.000 una baja. Y otra vez, ocurra lo que ocurra, habrá enviado dos predicciones correctas consecutivas a 16.000 personas. Iterando este procedimiento de concentrarse exclusivamente en la lista reducida de personas que han recibido sólo predicciones acertadas, puede crear en ellos la ilusión de que sabe de qué va la cosa. Al fin y al cabo, las 1.000 personas que habrán recibido 6 predicciones acertadas y ninguna equivocada (por coincidencia) tienen buenos motivos para desembolsar los 1000 dólares que les pide el editor del boletín: quieren seguir recibiendo estas declaraciones «proféticas».

Repito que una cuestión importante que hay que tener en cuenta al hablar de las coincidencias es la distinción entre clases genéricas de sucesos y

sucesos concretos. En muchas ocasiones la realización de un suceso particular es algo bastante raro —que a determinada persona le toque la lotería o que me llegue una determinada mano de bridge— mientras que el resultado genérico —que a alguien le toque la lotería o que salga esa mano de bridge— no tiene nada de extraordinario. Volvamos al problema del cumpleaños. Si sólo pedimos que 2 personas cumplan años el mismo día sin precisar cuál es este día particular, entonces bastan 23 personas para que el suceso tenga una probabilidad de $1/2$. Por contra, hacen falta 253 personas para tener probabilidad $1/2$ de que una de ellas tenga una fecha de cumpleaños determinada, el 4 de julio, pongamos por caso. Los sucesos concretos explicitados de antemano son, por supuesto, muy difíciles de predecir. Así pues, no es sorprendente que las predicciones de los televangelistas, curanderos, etc. suelen ser vagas y amorfas (hasta que se han producido los sucesos en cuestión, claro está, pues en ese instante los pronosticadores suelen afirmar que precisamente esos resultados son los que habían predicho).

Esto me recuerda el llamado efecto Jeane Dixon, por el cual las pocas predicciones acertadas (ya sea de los psíquicos, de los boletines bursátiles de pacotilla o de quien sea) se anuncian a bombo y platillo, mientras que las aproximadamente 9.800 predicciones fallidas hechas anualmente son oportunamente ignoradas. Se trata de un fenómeno muy general que contribuye a la tendencia que tenemos todos a dar a las coincidencias más importancia de la que en realidad merecen. Nos olvidamos de todas las premoniciones fallidas de desastres que hayamos tenido y recordamos vívidamente las que parecen acertadas. Cualquiera de nuestros conocidos ha oído hablar de ejemplos de telepatía; el número incomparablemente mayor de veces en los que no se ha producido es demasiado banal para ser tenido en cuenta.

Hasta nuestra biología parece conspirar para que las coincidencias parezcan más significativas de lo que realmente son. Como el mundo natural de rocas, plantas y ríos no parece ofrecer muchas pruebas de coincidencias superfluas, el hombre primitivo tenía que ser muy sensible a todas las anomalías y sucesos improbables imaginables a medida que iba construyendo la ciencia y su progenitor, el «sentido común». Al fin y al cabo, las

coincidencias «son» a veces muy importantes y significativas. Sin embargo, en nuestro complicado y, en gran parte, artificial mundo de hoy, la plétora de relaciones entre nosotros parece haber sobreestimulado la tendencia innata de mucha gente a notar la coincidencia y lo improbable, y les lleva a postular causas y fuerzas allí donde no hay nada. La gente conoce más nombres (además de los de los familiares, los de compañeros de trabajo y la gente famosa), fechas (desde artículos periodísticos hasta citas personales y programas), direcciones (ya sea de direcciones reales o números de teléfono, números de despacho, etc.) y organizaciones y acrónimos (desde el FBI al IMF, del SIDA al ASEAN) que en ningún otro momento del pasado. Por tanto, aunque sea muy difícil de cuantificar, el ritmo al que se producen las coincidencias probablemente ha aumentado en el último siglo. Y, a pesar de todo, no tiene mucho sentido buscar una explicación a la mayoría de ellas.

En realidad, la coincidencia más asombrosamente increíble que se pueda imaginar es la falta absoluta de coincidencias.

[Breves deducciones de los enunciados del cumpleaños (véase también la entrada sobre *Probabilidad*): (1) La probabilidad de que 2 personas tengan distinto cumpleaños es $364/365$; la de que 3 personas tengan distinto cumpleaños es $(364/365 \times 363/365)$; la de que 4 personas $(364/365 \times 363/365 \times 362/365)$; la de que 23 $(364/365 \times 363/365 \times 362/365 \times \dots \times 344/365 \times 343/365)$, producto que resulta ser igual a $1/2$. Por tanto, la probabilidad complementaria de que por lo menos dos personas tengan el mismo cumpleaños es también $1/2$ (1 menos el producto anterior). (2) La probabilidad de que alguien no cumpla los años el 4 de julio es $364/365$; la probabilidad de que de 2 personas ninguna cumpla años el 4 de julio es $(363/365)^2$; la probabilidad con 3 personas es $(363/365)^3$, y con 253 personas es $(363/365)^{253}$, que resulta ser $1/2$. Por tanto, la probabilidad complementaria de que al menos una de las 253 personas cumpla años el 4 de julio es también $1/2$, $1 - (363/365)^{253}$].

Combinatoria, grafos y mapas

Imagine que usted es miembro de un gremio de artesanos en extinción conocido como los coloreadores de mapas, y tiene ante sí un mapa plano de las Reticuladas Regiones de Convolúcido. El servicio de publicaciones de la universidad que le ha contratado pasa una mala época y usted se pregunta si podrá colorear el mapa con, a lo sumo, cuatro colores, asegurando, por supuesto, que los países vecinos que tengan un trozo de frontera común estén pintados en distinto color. El teorema de los cuatro colores garantiza que usted podrá cumplir con esta tarea sin importar cuantos países haya ni como estén dispuestos, siempre y cuando sean regiones conexas (esto es, no puede haber un pedazo de Esquizostán aquí y otro mil kilómetros más allá).

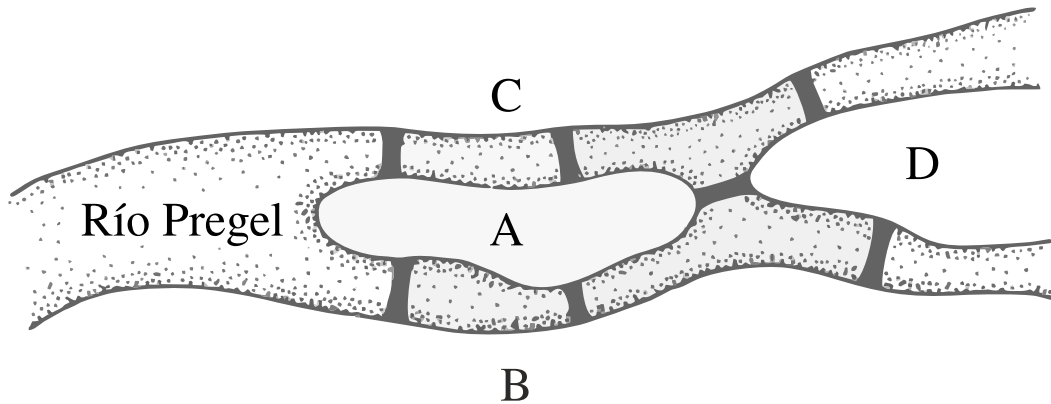
La conjetura de que para un mapa plano basta con cuatro colores fue formulada a mediados del siglo XIX, pero a pesar del interés que se puso en ello, quedó sin resolver hasta 1976, cuando los matemáticos norteamericanos Kenneth Appel y Wolfgang Haken demostraron que bastaban efectivamente cuatro colores. Anteriormente se había podido demostrar que como máximo hacían falta siete colores para pintar mapas sobre la superficie de un toro (una figura en forma de cámara de neumático), pero como la gente no suele dibujar mapas sobre roscones, este resultado carecía del atractivo natural del teorema de los cuatro colores.

Para demostrar el teorema de los cuatro colores hace falta un ordenador para poder examinar la miríada de posibilidades asociadas a los diversos tipos de configuraciones de mapas. Se trata de un nuevo progreso que, a primera vista, está reñido con la idea misma de demostración matemática. Tal y como se ha entendido tácitamente desde el tiempo de los griegos, las demostraciones han de ser comprensibles y verificables por los humanos.

Además han de ser lógicamente convincentes. (Véase la entrada sobre *QED*). Una demostración como la del teorema de los cuatro colores, que precisa de un uso tan generalizado de ordenadores, no es comprensible ni verificable en el mismo sentido que otras demostraciones matemáticas. Ni tampoco es lógicamente cerrada. Aunque la probabilidad de error en una configuración dada sea mínima, el número de permutaciones y ordenaciones que hay que examinar es tan inmenso que sólo podemos concluir que el teorema es probablemente verdadero. Pero «probablemente verdadero» no es lo mismo que «demostrado concluyentemente».

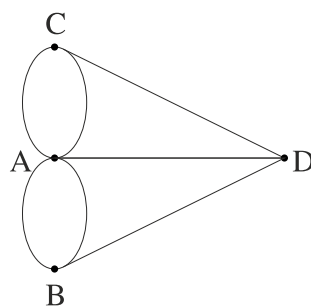
Algunos matemáticos han observado que en la teoría de grupos hay algunos teoremas tan complicados y largos que se les podrían plantear las mismas objeciones aunque no hagan falta ordenadores para demostrarlos. Como quiera que uno considere este asunto, resulta claro que la demostración del teorema de los cuatro colores no es elegante, convincente ni natural. Ciertamente, no forma parte de lo que el matemático Paul Erdős llama El Libro de Dios, un conjunto de demostraciones ideales que tienen estas propiedades. A pesar de todo, es una solución impresionante de un viejo problema.

Desgraciadamente, las secuelas matemáticas que ha traído el resultado no han sido tan impresionantes como las de otro viejo enigma planteado por Leonhard Euler, el problema de los puentes de Königsberg. Euler empezó su clásico artículo de 1736 (que muchos toman como punto de partida de la combinatoria) hablando acerca de la distribución de la ciudad de Königsberg, en Prusia oriental. La ciudad se asienta en las riberas del río Pregel, con dos islas en medio. Las diversas partes de la ciudad estaban unidas por siete puentes y los domingos la gente solía pasearse por ellas. La cuestión que se planteaba era si los habitantes de esa ciudad podían salir de su casa y volver a ella después del paseo habiendo atravesado cada puente del río una sola vez. Euler demostró que tal ruta no existía. La idea fundamental de la demostración consistía en que el número de entradas de dicha ruta a cada parte de la ciudad tendría que ser igual al número de salidas, lo cual implicaría que cada parte de la ciudad ha de tener un número par de puentes, condición que no se daba en Königsberg.



Los siete puentes de Königsberg

Euler representó las distintas partes de la ciudad por puntos y los puentes que las unían por líneas. El conjunto de puntos y líneas (o de vértices y aristas) resultante forman un grafo, y en el resto de su artículo Euler estudió el problema general siguiente: ¿en qué grafos de este tipo es posible encontrar un camino que pase sólo una vez por cada línea hasta regresar al punto de partida? (Obsérvese que, contrariamente a lo que sucede en el caso de Königsberg, es posible encontrar dicho camino para un grafo en Estrella de David, construido por superposición de dos triángulos, uno apuntando hacia arriba y otro un poco más caído apuntando hacia abajo. Nótese también que de cada vértice de este último grafo parte un número par de aristas).



Grafo de la ciudad según Euler

Esta manera aparentemente ingenua de representar relaciones matemáticas por medio de puntos y líneas que los unen es actualmente un útil indispensable en la teoría de grafos. Puede servir, por ejemplo, para representar la red de relaciones entre un grupo de personas (las líneas unen a

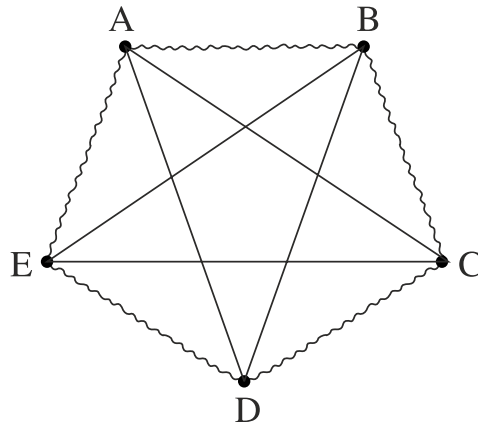
los que se conocen), el árbol de todos los posibles resultados de una sucesión de elecciones, los emparejamientos en rueda de un torneo deportivo, las conexiones de un circuito integrado y muchas otras cosas. Por contra, el teorema de los cuatro colores parece ser un callejón sin salida.

Es imposible predecir si un problema traerá nuevos progresos y sugerirá nuevas ideas o si llevará a una vía muerta. ¿Cuál de los dos acertijos siguientes es un callejón sin salida?

Tómese un entero positivo y, si es par, divídase por 2, pero si es impar multiplíquese por 3 y añádasele 1. Aplíquese la misma regla al entero resultante e itérese el proceso. La sucesión generada a partir de 11 es: 11, 34, 17, 52, 26, 13, 40, 20, 10, 5, 16, 8, 4, 2, 1, 4, 2, 1, ..., mientras que la generada por 92 es: 92, 46, 23, 70, 35, 106, 53, 160, 80, 40, 20, 10, 5, 16, 8, 4, 2, 1, 4, 2, 1, ... La cuestión es si cualquier número positivo acaba por caer en el ciclo 4-2-1 y, aunque se cree que esto es cierto, nadie ha podido demostrarlo todavía.

El segundo problema se puede formular en términos de invitados a un pequeño banquete. La cuestión es la siguiente: ¿Cuál es el menor número de comensales necesario para asegurar que al menos 3 de ellos se conozcan entre sí o al menos 3 de ellos no se conozcan? (Se supone que si Marta conoce a Jorge, entonces Jorge conoce a Marta). Se puede ver fácilmente que la respuesta es 6 tomando uno de los invitados al banquete, Juan. Como conoce o no conoce a cada uno de los otros 5 invitados, seguro que conoce por lo menos a 3 de ellos o no conoce por lo menos a 3 de ellos. Supongamos que Juan conoce a 3 comensales (en el caso de que no conozca por lo menos a 3 el razonamiento sigue un camino análogo) y pensemos en las posibles relaciones entre ellos. Si algún par entre ellos se conocen, éstos y Juan forman un grupo de 3 comensales que se conocen entre sí. Y por otra parte, si ninguno de los 3 conocidos de Juan conoce a los otros dos, ellos mismos ya forman un grupo de 3 comensales que no se conocen. Por tanto 6 son suficientes. Para ver que 5 comensales no bastan, imaginemos otra vez a Juan en una reunión de esta clase donde conoce exactamente a 2 de los otros 4 comensales y que cada uno de ellos conoce a una persona distinta de las que Juan no conoce.

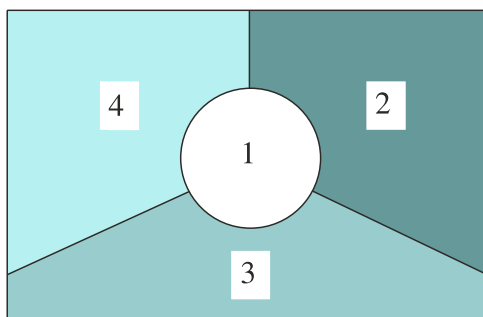
P ————— Q – P conoce a Q
 R ~~~~~ S – R no conoce a S



No hay ningún conjunto de 3 personas que se conozcan entre sí, ni tampoco un conjunto de 3 personas que no se conozcan

En un metanivel la pregunta es: ¿cuál de las dos preguntas conduce a algo más y cuál no? La respuesta es que, aunque desde un punto de vista práctico ninguna de las dos sirve de gran cosa, la primera es un callejón sin salida (al menos por ahora), mientras que la segunda lleva a una rama completamente nueva de la combinatoria, la teoría de Ramsey. Recibe este nombre por el matemático inglés Frank Ramsey y se ocupa de encontrar el mínimo número de elementos necesarios que satisfacen varias condiciones combinatorias simples. Hay un sinnúmero de problemas que son terriblemente fáciles de enunciar y aún no están resueltos, siquiera por la fuerza bruta de los métodos informáticos. Uno de tales problemas es la generalización del problema anterior que pregunta el mínimo número de comensales necesarios para garantizar que haya al menos 5 invitados que se conozcan entre sí o al menos 5 que no se conozcan.

Un ejercicio final. El teorema de los cuatro colores dice que como máximo hacen falta cuatro colores para pintar cualquier mapa plano. Hay varios mapas pequeños esencialmente distintos que muestran que por lo menos hacen falta cuatro colores. Dibuje uno.



Mapa que precisa de cuatro colores. Cuatro colores son siempre suficientes

La complejidad de un programa

Una vez conocí a una mujer que había leído un libro sobre recursos mnemotécnicos y, según parece, lo había malinterpretado completamente. Para memorizar un número de teléfono, por ejemplo, quizá tenía que recordar que su mejor amiga tenía 2 hijos, su dentista tenía 5, su compañera de piso en la universidad 3, su vecino de la derecha tenía 3 perros, el de la izquierda 7 gatos, su hermano mayor tenía 8 hijos si se contaban los de todas sus esposas y que en su casa eran 4 hermanos. El número de teléfono sería 253 37 84. Sus algoritmos (recetas, programas) eran enrevesados, ingeniosos, divertidos y siempre muchísimo más largos que lo que le pretendían ayudar a recordar. A veces, naturalmente, cuando tales algoritmos están íntimamente relacionados con una historia o un episodio que uno conoce muy bien, su longitud es sólo aparente y es razonable y está justificado usarlos. Sin embargo, no era éste el caso de mi amiga, que invariablemente acababa olvidando algún elemento esencial.

En el supuesto de que usted tuviera algún interés en hacerlo, ¿cómo describiría las sucesiones siguientes a un conocido que no pudiera verlas?

(1) 0 0 1 0 0 1 0 0 1 0 0 1 0 0 1 0 0 1 0 0 1 0 ...

(2) 0 1 0 1 1 0 1 0 1 0 1 0 1 0 1 1 0 1 0 1 0 1 1 ...

(3) 1 0 0 0 1 0 1 1 0 1 1 0 1 1 0 0 0 1 0 1 0 1 1 0 0 ...

La primera sucesión es claramente la más sencilla, pues es una simple repetición de dos ceros y un 1. La segunda sucesión presenta alguna irregularidad —un solo 0 alternando a veces con un 1 y a veces con dos unos— mientras que la tercera sucesión es la más difícil de describir, pues no da

señales de seguir ningún patrón. Nótese que el significado de los puntos suspensivos de la primera sucesión es evidente; lo es menos en la segunda, y no lo es en absoluto en la tercera. A pesar de esto, supongamos que cada una de estas sucesiones tiene un billón de bits de longitud (un bit es un 0 o un 1) y continúan «del mismo modo».

Atentos a estos ejemplos, podemos seguir al informático Gregory Chaitin y al matemático ruso A. N. Kolmogorov, y definir la complejidad de una sucesión de ceros y unos como la longitud del programa de ordenador más corto que genere dicha sucesión (esto es, que la imprima).

¿Qué quiere decir esto? Nótese que un programa que imprima la primera sucesión puede consistir simplemente en la siguiente receta: «Imprimir dos ceros, luego un 1, y repetir un tercio de billón de veces». Dicho programa es muy corto, especialmente si se compara con la longitud de la sucesión de un billón de bits que genera. La complejidad de esta primera sucesión puede ser de sólo unos 1000 bits, o la longitud que tenga el programa más corto que la produzca. Hasta cierto punto esto depende del lenguaje empleado para escribir el programa, pero independientemente de cuál sea ese lenguaje, al traducir a él «Imprimir dos ceros y luego un 1» no nos dará un programa muy largo. (Para uniformizar podemos suponer también que nuestros programas están escritos en un lenguaje de máquina muy sobrio consistente en 0 y 1, de modo que los programas que generan las sucesiones son también a su vez sucesiones de 0 y 1).

Un programa para producir la segunda sucesión podría ser una traducción de lo siguiente: «Imprimir un 0 seguido de un solo 1 o dos unos, y la pauta de los unos intermedios es 12111112 etc.». Si esta pauta persiste, cualquier programa que imprima la sucesión habrá de ser muy largo, pues tendrá que especificar completamente la parte «etc.» de la pauta de los unos intermedios. Sin embargo, debido a la alternancia regular de ceros y unos, el programa será considerablemente más corto que la sucesión de un billón de bits que produce. Así, la complejidad de esta segunda sucesión podría ser de sólo medio billón de bits, o la longitud que tenga el programa más corto que la produzca.

Con la tercera sucesión (con mucho, el tipo más frecuente) la situación es distinta. La sucesión, lo supondremos así, es tan desordenada a lo largo de

todo el billón de bits que ningún programa que la pueda generar es más corto que la sucesión misma. Todo lo que un programa puede hacer en este caso es enumerar tontamente los bits de la sucesión: «Imprimir 0, luego 1, luego 0, luego 0, luego 1, luego 0, luego 1, ...». No hay modo de comprimir los puntos suspensivos ni de acortar el programa. Dicho programa sería como mínimo tan largo como la propia sucesión que ha de imprimir, con lo que la tercera sucesión tiene una complejidad de aproximadamente un billón de bits.

Una sucesión como la tercera, que precisa de un programa tan largo como ella misma para reproducirla, se llama aleatoria. Las sucesiones aleatorias no muestran ninguna pauta, regularidad ni orden, y los programas que las generan no pueden hacer sino copiarlas directamente: «Imprimir 10001011011 ...». Estos programas no se pueden condensar ni abreviar. La complejidad de las sucesiones que producen es igual a la longitud de las mismas. Por contra, las sucesiones ordenadas y regulares como la primera se pueden generar con programas muy cortos y su complejidad es mucho menor que su longitud.

En ciertos aspectos, las sucesiones como la segunda son las más interesantes, pues, al igual que los seres vivos, presentan elementos de orden y elementos de azar. Su complejidad es menor que su longitud, pero no es tan pequeña como si fuera completamente ordenada ni tan grande como si fuera aleatoria. En lo que respecta a su regularidad, la primera sucesión de los ejemplos presentados se podría comparar a un diamante o a un cristal de cuarzo, mientras que la tercera sería comparable, por su aleatoriedad, a una nube de gas o a una sucesión de lanzamientos de una moneda. El análogo de la segunda podría ser algo así como una rosa (o una alcachofa, para no ser tan poéticos), que presenta a la vez orden y azar entre sus partes.

Estas comparaciones son algo más que simples metáforas. La razón de ello es que la mayoría de fenómenos se pueden describir por un código que se puede digitalizar y reducir a sucesiones de ceros y unos, tanto si es el lenguaje molecular de los aminoácidos y las moléculas de ADN como si es el idioma español de las cartas y los libros. (Véanse también las entradas sobre *Música, pintura y digitalización* y *Test de Turing*). Tanto el ADN como una novela romántica son, en sus respectivos códigos, sucesiones como la del segundo ejemplo, y presentan tanto orden y redundancia como azar y

complejidad. De modo análogo, las melodías complejas se encuentran entre las simples repeticiones de golpes y el ruido informe (análogos respectivamente a las sucesiones como la primera y la tercera).

El conjunto de la ciencia podría concebirse bajo este prisma. Ray J. Solomonoff y otros autores han teorizado que las observaciones de un científico podrían codificarse en una sucesión de ceros y unos. El objeto de la ciencia sería entonces encontrar programas cortos capaces de generar dichas observaciones (derivarlas o predecirlas). Un programa semejante, según dicen, sería una teoría científica, tanto más potente cuanto más corta fuera con respecto a los fenómenos experimentales predichos. Los sucesos aleatorios no serían predictibles, excepto en el sentido muy pickwickiano de un programa que simplemente los enumerara. Tal concepción de la ciencia es bastante simplista y sólo empieza a cobrar sentido cuando se refiere a un marco científico bien definido y fijado de antemano. Sin embargo, hace pensar en la generalidad de la idea de complejidad.

El concepto de complejidad que hemos definido aquí se llama complejidad algorítmica porque se mide por la longitud del programa (algoritmo o receta) más corto necesario para generar una sucesión dada. Un concepto afín es el de complejidad computacional de una sucesión, y se define como el tiempo más corto empleado por un programa que la genere. En los problemas de aplicación práctica el tiempo es a menudo el factor más importante que determina la elección de un programa. Quizás un programa corto tarde eones en generar la sucesión que se quiere (o la predicción, según las ideas de Solomonoff sobre la ciencia), mientras que un programa más largo lo haga en menos tiempo.

La mayoría de problemas de cálculo (y recordemos que, al igual que sus soluciones, siempre se pueden codificar en sucesiones) tardan más tiempo en ser resueltos cuanto mayores son. Por tanto, el conocido problema del viajante, que pide encontrar la ruta más corta que pasa una vez por cada una de las ciudades de una cierta región y que regresa al punto de partida, es tanto más largo de resolver cuantas más ciudades tengamos que considerar.

Es importante distinguir entre los problemas matemáticos cuya complejidad computacional es una función exponencial de su tamaño y aquéllos en los cuales es una potencia del tamaño. Como ilustración de esto,

en el caso del problema del viajante el tamaño está determinado por N , el número de ciudades a visitar, y no se sabe con certeza si el tiempo necesario para resolver el problema crece exponencialmente como 2^N o como una potencia N^T , donde T es un cierto exponente (generalmente pequeño). (Véase la entrada sobre *Crecimiento exponencial*). Un ejemplo numérico resultará aquí ilustrativo. Si tomamos $N = 20$ y $T = 3$, encontramos que $2^N = 2^{20} = 1.048.576$ mientras que $N^T = 20^3 = 8.000$ (solamente).

Si, como se sospecha, la complejidad computacional del problema del viajante crece como 2^N , entonces el tiempo necesario para resolverlo aumenta tan rápidamente que, a efectos prácticos, el problema es irresoluble para valores muy grandes de N . (Esto no cambiaría tampoco con los procesadores en paralelo, ordenadores cuya arquitectura interna les permite realizar muchas operaciones simultáneamente en paralelo, en contraposición a los procesadores en serie de los ordenadores de hoy en día, que realizan una operación cada vez). Si la complejidad computacional crece polinómicamente como una cierta potencia de N , entonces se pueden tener soluciones completas al problema en un tiempo razonable, incluso para valores grandes de N .

La íntima alianza conceptual entre los conceptos de economía y condicionante y los de complejidad algorítmica y computacional explica la importancia creciente de éstos. La capacidad de almacenaje de cualquier ordenador es limitada y, para muchos problemas, hay que idear algoritmos más eficaces (los horarios e itinerarios de unas líneas aéreas, por ejemplo). Como ya he dicho, estos conceptos de complejidad son también importantes desde un punto de vista teórico y filosófico. Muchos teoremas, incluido el teorema de incompletitud de Gödel (véase la entrada sobre *Gödel*), tienen demostraciones muy naturales si se expresan en términos de complejidad. Otros enunciados, como el llamado problema de P-NP (que pregunta si los problemas cuyas soluciones propuestas se pueden comprobar rápidamente [polinómicamente] son tales que sus soluciones se pueden descubrir también rápidamente), sólo tienen sentido en el contexto de la teoría de la complejidad.

La conciencia humana y su naturaleza fractal

Los fractales son curvas, superficies o figuras geométricas de dimensión superior que tienen la propiedad de conservar su estructura característica al ampliarlas, de modo que siguen presentando el mismo tipo de complejidad si nos las miramos cada vez de más cerca. (Véase la entrada sobre *Fractales*). Esta autosemejanza nos hace pensar en la conciencia humana, que parece estar expresamente habituada a ella, tanto si uno está pensando lógicamente con algún objetivo concreto como si está meditando distraídamente sin rumbo fijo o si, cuando algo capta su interés, se detiene para echarle una mirada más detallada que, a su vez, puede remitirle a seguir explorando con mayor detalle o de vuelta a la línea de pensamiento original.

Las discusiones serias sobre la reducción internacional de armamento y las chanzas de barbería tienen en común una «forma» humana característica y, como ha sugerido el matemático Rudy Rucker, el movimiento de avance, la digresión horizontal, la ramificación y la vuelta atrás en los distintos niveles y escalas son trazos que definen esta forma humana y constituyen un fractal en un espacio lógico multidimensional. Dada la vaguedad de esta última frase, y como su tono recuerda al de Jorge Luis Borges, intentaré aclararla con un cuento al estilo de este escritor argentino (a cuyo protagonista daremos el epónimo de Rucker). El relato que sigue toma la forma de reseña de un libro imaginario. Comentar este libro que no ha sido escrito resulta muchísimo más fácil que escribirlo.

RUCKER: UNA VIDA FRACTAL, por Eli Halberstam

Editado por Belford Books, Boston, 3.213 págs.,
3.900 ptas.

Versión en disco de Peaches N' Cream Software,
Atlanta, 7.900 ptas.

Reseña de Paul John Allen

El célebre matemático Eli Halberstam (galardonado con la codiciada Medalla Fields en matemáticas y autor de los libros de éxito *Asuntos de la mente* y *Caos, elección y azar*) ha escrito una primera y colosal novela basada en el concepto esotérico de fractal inventado por Benoît Mandelbrot y no cabe la menor duda de que hasta el momento nunca se había intentado nada semejante —ni tan siquiera por autores como James Joyce o Marcel Proust—. No importa por donde se empiece el libro, pues más que leerlo hay que vagar y curiosear a lo largo de sus páginas. No hay un hilo conductor convencional, sino más bien una multitud indefinida de excursiones, unificadas todas por la conciencia de un tal Marvin Rucker, el otro yo de Halberstam.

El volumen de 3.213 páginas arranca en el estudio del profesor Rucker, un matemático de mediana edad, donde éste está intentando aclararse con algunos tediosos teoremas relacionados con el conocido problema de $NP = P$. La verdadera novedad, sin embargo, se explica en la introducción, donde se informa al lector (hojeador) que, después de leer un episodio, puede seguir linealmente hacia adelante, volver sobre sus pasos a un episodio anterior o moverse horizontalmente, concentrándose en cualquier palabra o frase importante del mismo episodio, dirigiéndose luego a una elaboración posterior sobre la misma. Aunque suene bastante simple, el movimiento se demuestra andando o, como el propio Rucker piensa para sí, un tanto prosaicamente, en varias ocasiones, «Dios está en los detalles».

Por ejemplo, Rucker se hurga las narices mientras está pensando en sus teoremas y, si el lector escoge investigar sobre esto hasta el fin, es enviado a una página (en la versión en disco las alternativas se presentan en un menú que aparece en la parte inferior del monitor) donde se discute detenidamente la afición de Rucker por las prospecciones proboscideas. ¿Qué porcentaje de gente se hurga las narices? ¿Por qué tan pocos lo hacen en público y por qué,

sin embargo, tantos se abandonan a este placer en la falsa intimidad de sus automóviles? Si avanza un poco más en esta dirección, hay el recuerdo de unas pocas semanas atrás cuando Rucker, parado en un semáforo, vio a la señorita Samaras elegantemente peinada, sentada en el BMW de enfrente, con el índice profundamente clavado, aparentemente en el córtex frontal.

Si se cansa de esto, puede retroceder y volver al estudio de Rucker, donde acaba de entrar su hijo menor, babeando chorretones de caramelo barbilla abajo. Rucker está a punto de reprenderle suavemente por estropear su nueva calculadora cuando recuerda cómo de joven le gustaba mascar caramelos blandos. Como antes, el lector puede seguir adelante con la historia o investigar a fondo sobre los caramelos infantiles, los padres preocupados o el tono de voz que uno usa para regañar a los niños. Cada alternativa nos remite también a varias otras. La gracia de esta proliferación arbórea es la sensación vital, evanescente y frágil que da al libro.

Halberstam aconseja al lector que lea sólo los relatos, apartes y viñetas que despierten su curiosidad; como máximo una cuarta parte del libro. La versión para ordenador tiene al final un pequeño cuestionario cuyas respuestas dependen de las partes del libro que haya seleccionado el lector. Algunos amigos y colegas míos leyeron independientemente el libro en pantalla de vídeo y, como en *Rashomon*, nuestras respuestas a las preguntas del cuestionario fueron sustancialmente diferentes, y en la forma predicha por el ordenador, que había registrado los pasajes que había elegido cada uno.

Ni siquiera un libro gigantesco como este basta para desarrollar todas las bifurcaciones que pueden tomar los distintos relatos. Pero la habilidad artística de Halberstam supera esta explosión combinatoria de posibilidades y liga y entreteje imperceptiblemente el material, creando la ilusión de una bifurcación ilimitada. Hay varios relatos principales: uno de ellos narra la complicada vida familiar de Rucker, otro trata de un juego de timo casi ilegal, y un tercero ilustra de una manera interesante las ideas más destacadas de la teoría de la complejidad, una nueva y apasionante rama de la informática y la lógica matemática.

Ocurre con frecuencia que en las coyunturas delicadas hay pocas alternativas, si las hay. El efecto que se pretende, como un río desbocado, es sugerir la simpleza del protagonista en esas ocasiones. Por ejemplo, medio

por curiosidad medio por lascivia, Rucker, un bufón intelectual que recuerda vagamente un personaje de Saul Bellow, ha marcado un teléfono erótico de la línea 903. Después de haber entrado en materia, oye un zumbido que indica que tiene otra llamada esperando. Presiona dos veces el receptor y descubre que es su mujer que le llama desde el supermercado. Aturrullado, trata de acabar pronto con ella y le dice que está hablando con un colega de la escuela, a lo que ella exclama que tiene que hablar con él personalmente sobre la faja que ha elegido para el próximo *bar mitzvah* de su hijo, y si por favor puede darse un poco de prisa y marcar la tecla «R» para que la llamada se convierta en una multiconferencia. Naturalmente no puede, pues la lúbrica dama de la otra línea añadiría una nueva mella a su ya maltrecha relación conyugal.

A pesar de esos giros narrativos, esta matriz casi sensible de desviación, digresión y movimiento horizontal en la obra es lo que vivifica a Rucker y sus hazañas, y lo que más impresiona al lector. Los detalles, grandes y pequeños, sobre temas críticos y banales, van sucediéndose en esta crónica barroca y multidimensional. A los que entendemos de matemáticas, Halberstam parece estar diciéndonos que la mejor manera de modelizar la consciencia humana —como las costas infinitamente melladas, las arrugadas y varicosas superficies de las montañas, los remolinos y espirales del agua turbulenta, o una multitud de otros fenómenos «fracturados»— es echar mano del concepto geométrico de fractal. La definición no es importante aquí, pero sus características más específicas son la ramificación y complejidad ilimitadas, así como la propiedad peculiar de la autosemejanza, por la cual un objeto fractal (el libro, en este caso) presenta el mismo aspecto independientemente de cuál sea la escala a la que es examinado (tanto los hechos principales como los detalles más finos).

En vez de seguir disertando sobre esto (Halberstam no lo hace), me contentaré con observar que, al manifestar la inagotable capacidad de divagación del hombre, el libro muestra también la unidad e integridad personal de la consciencia humana. La estructura de la obra es virtuosa y, aunque no es preciso que John Updike se inquiete, la obra es muy útil —casi todo lo que podría desearse dada su extensión—. El libro lleva con facilidad su carga didáctica y, a pesar de su extensión, uno se desprende de él habiendo

captado la idea vivida y precisa de una persona suplente: Marvin Rucker. Los episodios son densos; de hecho, no se pueden prácticamente sintetizar, y resumir más la trama sería engañosamente reduccionista. Sus parientes literarios más próximos son el *Ulysses* de Joyce y el *Tristram Shandy* de Laurence Sterne, pero los dos carecen de la musculatura cerebral de *Rucker: Una vida fractal*.

El libro es digno de un público amplio, al que, desgraciadamente, puede que no atraiga por el temor que le produce a mucha gente cualquier cosa que suene vagamente a matemáticas. Quizá sea despachado como una mera proeza técnica, mera ciencia ficción o un mero lo que sea, del que nuestros literatos, generalmente anuméricos, saben poco y por tanto mantienen una actitud de rechazo total. El hecho de que esta reseña tenga asignadas sólo 1.250 palabras es, en parte, una evidencia en favor de una opinión posiblemente paranoide como ésta.

Como el Herzog de Bellow, Rucker escribe cartas a un conjunto variado de personas, algunas famosas, otras no, algunas vivas y otras muertas. Uno de sus muchos «corresponsales» es Alexander Herzen, un escritor y disidente liberal ruso del siglo XIX. Rucker cita dos veces la famosa frase de Herzen, «El arte y el relámpago estival de la felicidad humana, éstos son nuestros únicos dones verdaderos». Quizás Halberstam se hacía eco de la yuxtaposición del relámpago, que tiene una estructura fractal, y el arte, que en este caso particular también la tiene. En cualquier caso, *Rucker: Una vida fractal* reparte los verdaderos dones.

Conjuntos infinitos

Lega usted al hotel, acalorado, sudoroso e impaciente. Su humor no mejora cuando el recepcionista le dice que no sabe nada de su reserva y que el hotel está lleno. «Me temo que no puedo hacer nada por usted», salmodia oficiosamente el empleado. Si usted tiene ganas de discutir, podría informar al recepcionista, en un tono igualmente oficioso, que el problema no es que el hotel esté lleno, sino que además de estar lleno es finito. Le podría explicar que si el hotel estuviera lleno pero fuera infinito, sí podría hacerse algo. Podría decir al individuo de la habitación 1 que fuera a la 2; que el individuo de ésta se fuera a la 3, cuyos ocupantes anteriores se habrían ido ya a la 4, etc. En general, los huéspedes de la habitación N se trasladarían a la $(N + 1)$ para cualquier número N . Con esta acción ningún individuo se quedaría sin habitación y la número 1 quedaría vacante para usted.

Los conjuntos infinitos tienen muchas propiedades sorprendentes que no tienen los finitos. Por definición, un conjunto infinito siempre puede ponerse en correspondencia biunívoca con uno de sus subconjuntos; esto es, tiene subconjuntos cuyos elementos se pueden emparejar uno a uno con los elementos del conjunto total. La escena anterior del hotel ilustra que si al conjunto de los números enteros positivos le quitamos el 1, quedan tantos números (2, 3, 4, 5, ...) como números enteros positivos hay en total (1, 2, 3, 4, ...). Del mismo modo, hay tantos números pares como números enteros, y tantos números enteros múltiplos de 17 como números enteros. El siguiente emparejamiento aclara esta última afirmación: 1, 17; 2, 34; 3, 51; 4, 68; 5, 85; 6, 102; etc.

Algunas de estas rarezas relativas a los conjuntos infinitos se conocían ya en tiempos de Galileo, pero su estudio sistemático lo debemos al matemático

alemán Georg Cantor, y el hecho de que la teoría de conjuntos se haya convertido en el lenguaje común de la matemática abstracta se debe en buena medida a sus resultados. No es pues sorprendente que el tema de los conjuntos infinitos sea muy extenso, y aquí sólo me entretendré en una distinción útil debida a Cantor: la diferencia entre los conjuntos infinitos numerables y los no numerables. Esta distinción tiene un papel muy importante en el análisis matemático, y las demostraciones relacionadas con estos conceptos son particularmente bellas.

Un conjunto infinito es numerable si hay alguna manera de asociar sus elementos o emparejarlos uno a uno (sin dejarse ninguno) con los números enteros positivos. Un conjunto infinito es no numerable si sus elementos no se pueden emparejar de ninguna manera con los enteros positivos. Los conjuntos infinitos que hemos mencionado por ahora son numerables, pero antes de dar un ejemplo de conjunto infinito no numerable, esbozaré una demostración debida a Cantor que prueba que el conjunto de los números racionales también es numerable, a pesar de su densidad y de su aparente plenitud. Esto es, hay exactamente tantos números enteros como fracciones. (Véase también la entrada sobre *Números racionales e irracionales*).

¿Cómo podemos emparejar los números racionales (fracciones) con los enteros positivos? No podemos relacionar sin más el $1/1$ con el 1, el $2/1$ con el 2, el $3/1$ con el 3 y así sucesivamente, pues nos estaríamos dejando la mayoría de números racionales. Otros intentos más sofisticados presentan problemas similares. Un truco que funciona es considerar primero los números racionales tales que su numerador y su denominador sumen 2. Sólo hay uno — $1/1$ — y le asociaremos el número 1. Luego consideramos los racionales cuyos numerador y denominador sumen 3. Hay dos — $1/2$ y $2/1$ — y les asociamos los enteros 2 y 3. A continuación consideramos los racionales cuyos numerador y denominador suman 4 — $1/3$, $2/2$ y $3/1$ — y asociamos el siguiente entero, 4, a $1/3$, luego el 5 a $3/1$, e ignoramos el $2/2$ (y haremos lo mismo con todas las fracciones que no sean irreducibles). En el estadio siguiente, las fracciones cuyos numerador y denominador sumen 5, asociamos el 6 a $1/4$, el 7 a $2/3$, el 8 a $3/2$ y el 9 a $4/1$.

1/1	2/1	3/1	4/1	5/1	6/1	...
1/2	2/2	3/2	4/2	5/2	6/2	...
1/3	2/3	3/3	4/3	5/3	6/3	...
1/4	2/4	3/4	4/4	5/4	6/4	...
1/5	2/5	3/5	4/5	5/5	6/5	...
1/6	2/6	3/6	4/6	5/6	6/6	...
⋮	⋮	⋮	⋮	⋮	⋮	

Demostración de que los números racionales son numerables

Seguimos así, considerando en cada estadio sólo los racionales cuyos numerador y denominador sumen N , ordenándolos de modo creciente y asociándoles los enteros subsiguientes. Finalmente, cada número racional queda emparejado con un entero y llegamos a la conclusión de que el conjunto de las fracciones es infinito numerable. (Quizá le apetezca comprobar que el entero 13 está asociado a la fracción $2/5$).

Y ahora, un conjunto no numerable. Cantor demostró que el conjunto de todos los números reales (todos los decimales) es más numeroso (más infinito) que el conjunto de los enteros o el de los racionales. O dicho con mayor precisión, que no hay ninguna manera de emparejar los números reales con los enteros (o los racionales) sin que queden reales por emparejar. La demostración estándar de esta propiedad es de tipo indirecto y empieza suponiendo que sí existe tal emparejamiento. Supongamos, sólo para concretar un poco, que el número 1 se empareja con el número real 4,56733951..., el 2 con 189,31299008..., el 3 con 0,33933337..., el 4 con 23,54379802..., el 5 con 0,98962415..., el 6 con 6.219,31218462..., etc. ¿Cómo podemos estar seguros de que, independientemente de cómo continúe esta lista infinita (o cualquier otra que podamos construir), siempre dejará fuera algunos números reales?

1	4 . <u>5</u> 6 7 3 3 9 5 1 ...
2	189 . 3 <u>1</u> 2 9 9 0 0 8 ...
3	. 3 3 <u>9</u> 3 3 3 3 7 ...
4	23 . 5 4 3 <u>7</u> 9 8 0 2 ...
5	. 9 8 9 6 <u>2</u> 4 1 5 ...
6	6,219 . 3 1 2 1 8 <u>4</u> 6 2 ...
⋮	⋮

El número que empieza 0,620835 no aparece en ningún lugar de la lista pues difiere del N-ésimo número de ésta al menos en la N-ésima cifra decimal. Los números reales no son numerables

Para responder a esto, consideremos el número real comprendido entre 0 y 1 cuyo N-ésimo lugar decimal esté ocupado por un dígito una unidad mayor que el dígito subrayado en el N-ésimo lugar decimal del N-ésimo número de la lista. (Quizá quiera usted releer esta última frase con un poco más de calma). Para la lista particular anterior, el número al que me refiero empezaría 0,620835..., pues el primer dígito, 6, es 1 más 5, 2 es 1 más 1, 0 es 1 más 9, etc. Este último número no está en la lista, pues por definición difiere del primer número de ésta al menos en la primera cifra decimal, del segundo al menos en la segunda cifra decimal, del tercero al menos en la tercera cifra, y difiere del N-ésimo número de la lista al menos en la N-ésima cifra decimal.

Y esto es todo. QED. Aquí acaba la demostración. No importa qué lista infinita de números reales nos presenten, con una técnica similar siempre podremos construir un número real que no esté en la lista. La conclusión es que no hay ninguna manera de emparejar los números reales con los enteros sin dejarse ninguno. El conjunto de los números reales es infinito y no numerable, y se dice que su cardinalidad (infinita) es superior a la del conjunto de los enteros y de los racionales.

[Hay también conjuntos con cardinalidades superiores a la de los números reales. Como ejemplos tenemos el conjunto de todos los subconjuntos de los números reales, o el conjunto de todas las funciones de números reales. De

hecho, existe toda una jerarquía de cardinalidades infinitas que empieza por $\aleph_0$, el símbolo de Cantor para la cardinalidad de los enteros ($\aleph$ es alef, la primera letra del alfabeto hebreo). Sin embargo, como ya dije, la distinción entre conjuntos numerables y no numerables es la única que cuenta para muchos matemáticos].

Cantor conjeturó que no había ningún subconjunto de los reales que fuera más numeroso que los enteros y menos numeroso que los propios números reales (y, en consecuencia, sugirió que su cardinalidad se denotara con el símbolo $\aleph_1$). Esta especulación es lo que posteriormente se ha llamado hipótesis del continuo y nunca nadie la ha demostrado. Una razón poderosa es que, como han demostrado Gödel y el matemático americano Paul Cohen, es independiente de los otros axiomas de la teoría de conjuntos; tanto la hipótesis del continuo como su negación son consistentes con la teoría de conjuntos tal y como la entendemos en la actualidad. Un nuevo axioma, que fuera verosímil, podría decidir la cuestión, pero a pesar de los intentos de muchos eminentes lógicos y expertos en teoría de conjuntos (por no hablar de un quijotesco intento por mi parte con «conjuntos genéricos», pues no encajaría en el margen de esta página), nadie lo ha encontrado todavía.

Volviendo a nuestro ejemplo del Hotel Infinito después de este viaje tan agotador, observamos que si cada una de la infinidad numerable de habitaciones del hotel tuviera una infinidad numerable de ventanas, el conjunto total de ventanas todavía sería numerable. Esto es, el hotel no tendría más ventanas que habitaciones. (La demostración es parecida a la de la numerabilidad del conjunto de los racionales). Finalmente, acabaremos con una observación contundente: olvidemos por un instante nuestras limitaciones físicas e imaginemos que esta infinidad de ventanas está numerada y que a las 11:59, las ventanas de la 1 a la 10 se rompen y que la ventana 1 es reparada. Medio minuto después, se rompen las ventanas de la 11 a la 20, mientras la 2 es reparada. Un cuarto de minuto más tarde se rompen las ventanas de la 21 a la 30 y la 3 es reparada. La progresión está clara: un octavo de minuto más tarde... La pregunta es: ¿cuántas de estas ventanas se han roto y han sido reparadas a las 12:00? Y la respuesta es que a las 12:00 se han roto y han sido reparadas todas.

Es hora de pedir la cuenta.

Correlación, intervalos y tests

Los niños con pies más grandes tienen mejor ortografía. En zonas del sur de Estados Unidos, los condados con mayores tasas de divorcio generalmente tienen menores tasas de mortalidad. Los países que añaden flúor al agua tienen tasas de cáncer más altas que aquellos que no lo hacen. ¿Hemos de estirar los pies de nuestros hijos? ¿Conviene que haya más artículos invitando a la «dolce vida» en *Penthouse* y *Cosmopolitan*? ¿Es una conjura la fluoración?

Aunque existan estudios que establecen todos estos resultados, las anteriores interpretaciones de los mismos sólo son posibles si uno no distingue entre correlación y causalidad. (Resulta interesante notar que el filósofo David Hume sostenía que, en principio, no hay diferencia entre ambos conceptos. Sin embargo, a pesar de algunas semejanzas superficiales, los temas que él consideraba eran completamente distintos a éstos). Aunque hay varias clases y medidas distintas de correlación estadística, todas ellas indican que dos o más cantidades están relacionadas de algún modo y en cierto grado, pero no necesariamente que una sea causa de la otra. A menudo, las variaciones en las dos cantidades correlacionadas son el resultado de un tercer factor.

Los extraños resultados anteriores se explican fácilmente del modo siguiente. Los niños que tienen los pies más grandes tienen mejor ortografía porque son mayores, siendo su mayor edad la causa de que tengan los pies más grandes y de que su ortografía sea mejor, aunque esto último no sea tan seguro. La edad es también un factor importante en el segundo ejemplo, pues las parejas mayores se divorciarán probablemente menos y se morirán con mayor probabilidad que las de condados con perfiles demográficos más

jóvenes. Y las naciones que añaden flúor al agua son en general más sanas y se preocupan más por su salud, con lo que un gran porcentaje de sus ciudadanos viven lo bastante para enfermar de cáncer que es, en buena medida, una enfermedad de gente mayor.

Para la mayoría, son menos importantes las definiciones de las medidas efectivas de correlación que la simple distinción anterior entre correlación y causa. Pero demasiado a menudo la gente se queda hipnotizada con los detalles técnicos de los coeficientes de correlación, las rectas de regresión y las curvas de máximo ajuste, y olvida echar una mirada atrás y meditar acerca de la lógica de la situación. El fenómeno me recuerda a las personas (yo soy una de ellas: véase el final de la entrada sobre *Teoría de juegos*) que se compran un nuevo ordenador o un nuevo procesador de textos para trabajar más aprisa y luego pierden una cantidad exorbitante de tiempo obsesionadas con los detalles del *software* e inventando programas «atajo», que tardan algunas horas en componer y cuya invocación sólo ahorra presionar tres o cuatro teclas.

«Parece como si todo el mundo los comprara. Todo el país se ha vuelto loco con estas cosas». «¿Cómo lo saben, si sólo hablaron con 1000 personas?». Como podrían sugerir estas dos citas contradictorias, la idea de muestra aleatoria es otro concepto estadístico simple cuya importancia no siempre se aprecia plenamente. Sin una de tales muestras, una multitud de testimonios personales y de frases «lo dice todo el mundo» quizá signifique muy poco, mientras que con una de ellas, un número sorprendentemente pequeño de encuestados puede ser concluyente. Basándose en observaciones realizadas con esa muestra, un intervalo de confianza es una banda numérica (que varía ligeramente de una muestra a otra) escogida de modo que contenga el verdadero valor desconocido de alguna característica de la población considerada, con una probabilidad especificada de antemano (normalmente el 95%). Así, si encuestamos una muestra aleatoria de 1.000 personas y el 43% están a favor de las ideas expuestas en la Constitución, entonces hay una probabilidad de aproximadamente el 95% de que el porcentaje de toda la población que está a favor de dichas ideas esté comprendido entre el 40% y el 46%, $43\% \pm 3\%$.

Aunque hay una serie de cuestiones técnicas relativas al cálculo e

interpretación de los intervalos de confianza, no hace falta conocerlas para comprender las ideas fundamentales, igual que ocurre en el caso de la correlación. De hecho, si uno estudia a fondo los detalles de la estimación de los intervalos de confianza puede pasarle por alto el alcance limitado del método. No se trata de que 1.000 personas no basten para darnos este intervalo de confianza de $\pm 3\%$. Antes bien lo que ocurre es que esa estimación es muy sensible al planteamiento dado al problema o a la formulación de la pregunta. Si en el ejemplo anterior las ideas se hubieran *identificado* como procedentes de la Constitución, las respuestas probablemente hubieran sido completamente distintas. Las creencias, actitudes e intenciones de los encuestados no permiten cambiar a la ligera la formulación de una pregunta por otra extensionalmente equivalente.

La comprobación de hipótesis estadísticas es otro concepto estadístico cuya comprensión no precisa conocer previamente el aparato técnico. Se hace una suposición, se diseña y se realiza un experimento para comprobarla, y luego se hacen algunos cálculos para ver si los resultados del experimento son suficientemente probables atendiendo a la suposición. Si no lo son se descarta la suposición y, a veces, se acepta provisionalmente una hipótesis alternativa. Así pues, la estadística sirve más para descartar proposiciones que para confirmarlas.

Al aplicar este procedimiento se pueden cometer dos tipos de errores: el error del tipo I consiste en rechazar una hipótesis verdadera y el de tipo II se produce cuando se acepta una hipótesis falsa. Ésta es una distinción útil en contextos menos cuantitativos. Por ejemplo, cuando se trata de desembolsos de fondos gubernamentales, el estereotipo de liberal procura evitar los errores del tipo I (que quien lo merece no reciba lo que le toca), mientras que el estereotipo de conservador está más preocupado por evitar los errores del tipo II (que quien no lo merece reciba más que lo que le toca). Si se trata ahora de castigar los delitos, el conservador de caricatura está más interesado en evitar los errores del tipo I (que el culpable no reciba su merecido), mientras que el liberal se apura en evitar los errores del tipo II (que el inocente reciba un castigo innecesario).

La FDA^[4] debe evaluar las probabilidades relativas de caer en errores del tipo I (no dar el visto bueno a un buen fármaco) del tipo II (aceptar un mal

fármaco). Los empresarios preocupados por el control de calidad han de contrapesar los errores del tipo I (rechazar una muestra con muy pocos artículos defectuosos) y los errores del tipo II (dar por buena una muestra con demasiados artículos defectuosos). En estas situaciones y en otras parecidas, la lógica de la comprobación estadística nos será de mucho provecho, aun cuando no dé cifras concretas como resultado.

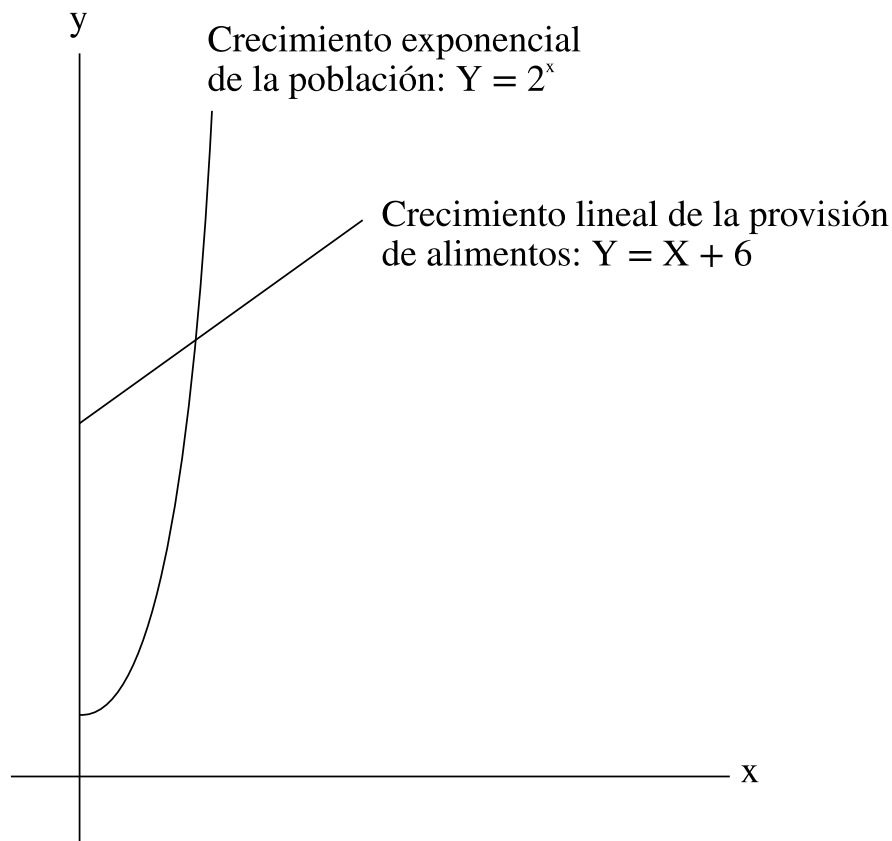
Más que la mayoría de las demás ramas de la matemática, la estadística es puro sentido común formalizado y pensamiento sencillo cuantificado. Las actitudes escépticas para con la estadística (como la de Benjamín Disraeli «Mentiras, malditas mentiras y estadística», o mi favorita, «Verdades, verdades a medias y estadística») están plenamente justificadas, pero no deberían hacernos perder de vista que se trata de una materia imprescindible. Renunciar a usarla sería cometer un error del tipo I (o, si uno tiene los pies pequeños, un error del tipo I).

Crecimiento exponencial

La sucesión de números 2, 4, 8, 16, 32, 64, ... crece exponencialmente, mientras que la sucesión 2, 4, 6, 8, 10, 12, ... lo hace linealmente. Tomando una adaptación del viejo cuento, observamos que si ponemos 2 céntimos en el primer escaque de un tablero de ajedrez, 4 en el segundo, 8 en el tercero, etc., el último escaque contendrá casi 200.000 billones de pesetas (2^{64} céntimos son aproximadamente $1,8 \times 10^{17}$ ptas.). Por contra, 2 céntimos en el primer escaque, 4 en el segundo, 6 en el tercero, etc., dan 1,28 pesetas en el último. En general, una sucesión crece exponencialmente (o también geoméricamente) si su tasa de aumento es proporcional a la cantidad presente, es decir, si cada número de la sucesión se obtiene multiplicando el anterior por un factor constante. Una sucesión crece linealmente (o también aritméticamente) si su tasa de aumento es constante, esto es, si cada término de la sucesión se obtiene añadiendo la misma cantidad a su antecesor.

Por supuesto, el factor por el que se pasa de un término de la sucesión al siguiente no ha de ser necesariamente dos. Por ejemplo, si las 1.000 ptas. que uno deposita hoy rinden un 10% anual, crecerán a razón de un factor 1,1 (o el 110%) cada año, y el próximo año valdrán 1.100 ptas. ($1000 \text{ ptas.} \times 1,1$). Al cabo de dos años valdrán 1.210 ptas. ($1.000 \text{ ptas.} \times 1,1 \times 1,1$), y al cabo de tres, 1.331 ptas. ($1.000 \text{ ptas.} \times 1,1 \times 1,1 \times 1,1$). Así pues, la sucesión 1.000 ptas., 1.100 ptas., 1.210 ptas., 1.331 ptas., ... es exponencial y al cabo de N años el valor del depósito original es de $1.000 \text{ ptas.} \times 1,1^N$. Si ese mismo dinero se hubiera depositado al 10% de interés simple, el valor del depósito original en los tres primeros años habría sido de 1.100 ptas., 1.200 ptas., 1.300 ptas., y al cabo de N años valdría $1.000 + 1.000 \times (0,10)N$ o, lo que es

lo mismo, $1.000 + 100N$.



Razonamiento de Malthus

Al igual que el dinero en un depósito a interés compuesto, las poblaciones (ya sean de personas o de bacterias) tienden a crecer exponencialmente, mientras que la producción de alimentos, al igual que el dinero puesto a interés simple, tiende a crecer sólo linealmente. De considerar ambas observaciones, el economista británico de principios del siglo XIX, Thomas Malthus, llegó a la conclusión de que la miseria y el hambre eran inevitables. Aunque su razonamiento tiene defectos y es atacable en varios puntos, su modo claro de articular la situación es admirablemente ilustrativo de la diferencia entre los crecimientos lineal y exponencial.

El crecimiento exponencial no sólo supera enseguida al lineal, sino que a la larga deja atrás al crecimiento cuadrático, al cúbico y a cualquier crecimiento polinómico en general. Una nueva mirada a las sucesiones con

las que hemos empezado nos servirá para aclarar esta jerga matemática. La sucesión 2, 4, 8, 16, 32, 64, 128, ... crece exponencialmente, y su término N-ésimo es 2^N . La sucesión 2, 4, 6, 8, 10, 12, 14, ... presenta un crecimiento lineal, y su término N-ésimo es igual a $2N$. Consideremos ahora la sucesión 1, 4, 9, 16, 25, 36, 49, ... Presenta un crecimiento cuadrático («cuadr» indica que los términos son cuadrados) y su término N-ésimo es igual a N^2 . El término N-ésimo de la sucesión cúbica 1, 8, 27, 64, 125, 216, ... es N^3 , mientras que el término N-ésimo de la sucesión 7, 42, 177, 532, ... es $2N^4 + 5N$.

La sucesión exponencial 2^N acaba por crecer más rápidamente que cualquiera de estas otras sucesiones, más rápidamente, de hecho, que cualquier sucesión polinómica (una que crezca como las potencias de N : N^2 , N^3 , N^4 , N^5 , etc.). Obsérvese que si, por ejemplo, $N = 30$, 2^N es 1.073.741.824, mientras que N^4 es sólo 810.000. Es particularmente importante evitar los crecimientos exponenciales en el diseño de métodos informáticos. (Véanse también las entradas sobre *La complejidad y Clasificar*). Los métodos que implican un tiempo de resolución que crece exponencialmente con el tamaño del problema generalmente tardan demasiado en ser resueltos y carecen de utilidad práctica. Si el problema tiene una gran cantidad de datos o es grande en algún otro sentido, podríamos tardar milenios en obtener la respuesta. Ni tan siquiera la velocidad de los superordenadores puede con el crecimiento exponencial. Por contra, los métodos cuyo tiempo de resolución crece linealmente, o como mucho polinómicamente, con el tamaño del problema se resuelven normalmente con suficiente rapidez como para que su aplicación sea de utilidad.

Ahora cambiaré de sentido y acabaré este breve comentario sobre el crecimiento exponencial con una variación en la que la sucesión disminuye exponencialmente. Cada término de una tal sucesión se obtiene multiplicando su antecesor por un número menor que 1. Así la sucesión 1000, 800, 640, 512, 409,6,... presenta lo que se llama decrecimiento exponencial porque cada término es el 80% del anterior (siendo el término N-ésimo después del primero igual a $1.000 \times 0,8^N$).

Un ejemplo importante de decaimiento exponencial lo tenemos en la desintegración de elementos radiactivos. A partir de la tasa de desintegración de dichos elementos se puede determinar su período de semidesintegración (el intervalo de tiempo que ha de transcurrir para que se desintegre la mitad de la sustancia). La datación por el carbono 14 se basa en cálculos de este tipo. La idea en la que se fundamenta este método consiste en que todos los seres vivos contienen una concentración conocida de carbono radiactivo, que al morir se desintegra. Midiendo cuánto carbono radiactivo queda, se puede calcular la edad de un carbón vegetal, de un vestido o de lo que sea.

Naturalmente no hace falta que nos traslademos a reinos tan arcanos para encontramos con el decaimiento exponencial. El ritmo al que pierde sabor el chicle de fruta también es exponencial. Otro ejemplo aún más próximo a casa —la mía—: se dice que por cada ecuación que uno añade a un libro de divulgación matemática o científica, se reduce a la mitad el número de potenciales lectores. Es decir, que con $A = B$ y $X = Y$ he eliminado al 75% de los lectores que, de otro modo, hubieran seguido más allá. La advertencia se podría formular también como, «El número de lectores decrece exponencialmente con el número de ecuaciones». Espero que no sea verdad.

Los cuantificadores en lógica

¿Es lo mismo engañar a todos un momento que engañar a algunos todo el rato? Para resolver esta punzante cuestión, nos concentraremos en los significados de «todos» y «algunos». La lógica proposicional elemental (véase la entrada sobre *Tautologías*) se puede considerar como el estudio completo de ciertas expresiones lógicas básicas: «si..., entonces...», «si y sólo si», «y», «o» y «no». Si añadimos a la lista «todo», «existe», «todos» y «algún», tenemos lo que se conoce como lógica predicativa, un sistema lógico que permite formalizar todo el razonamiento matemático. Estas últimas expresiones son lo que generalmente se llama cuantificadores (aunque quizá debieran llamarse cualificadores, pues no cuantifican mucho que digamos) y se usan para transformar formas relacionales en proposiciones declarativas. Así, la forma «X es calvo» se puede cuantificar universalmente de modo que diga «Todo X es calvo» o particularmente como «Algún X es calvo».

Consideremos la forma misantrópica «X odia a Y» en la que cada variable se puede cuantificar independientemente. Si ambas se cuantifican universalmente, se convierte en «Para todo X, para todo Y, X odia a Y», o dicho de un modo más natural «Todo el mundo odia a todo el mundo». Si la primera variable se cuantifica universalmente y la segunda existencialmente, tenemos «Para todo X, existe un Y (tal que) X odia a Y», o «Todo el mundo odia a alguien». Si cambiamos el orden de los cuantificadores en la frase anterior, obtenemos «Existe un Y (tal que) para todo X, X odia a Y» o, más coloquialmente, «Hay alguien odiado por todos». Si cuantificamos existencialmente la primera variable y universalmente la segunda, el resultado es «Existe un X (tal que) para todo Y, X odia a Y», que en un

español algo más claro reza «Hay alguien que odia a todo el mundo (incluso a sí mismo)». Y si ambas variables se cuantifican existencialmente, tenemos «Existe un X, existe un Y (tales que) X odia a Y», o «Alguien odia a alguien».

(El símbolo de «para todo» o «todo» es una A invertida y el de «existe» o «algún», una E al revés. Así, si simbolizamos «X odia a Y» por «H(X, Y)», entonces «Todo el mundo odia a alguien» se escribe simbólicamente « $\forall X \exists Y H(X, Y)$ », mientras que «Hay alguien odiado por todos» y «Hay alguien que se odia a sí mismo y a todos los demás» se convierten en « $\exists Y \forall X H(X, Y)$ » y « $\exists X \forall Y H(X, Y)$ », respectivamente. ¿Cómo se escribiría «Todo el mundo es odiado por alguien»?

La manipulación formal de los cuantificadores es especialmente importante en matemáticas donde podemos encontrarnos con una hilera de varios de ellos y es muy distinto «Para todo X existe un Y tal que para todo Z...» de «Existe un Z tal que para todo X hay un Y...». Un cuantificador mal colocado puede convertir una función continua en uniformemente continua, si usted sabe lo que quiero decir (y si no lo sabe también). Esos malabarismos son algo menos importantes en la vida cotidiana, pues el conocimiento del contexto de una frase permite evitar que se propaguen los errores. Es improbable, por ejemplo, que «Para cada X, existe un Y tal que Y es la madre de X» se confunda con «Para todo Y, existe un X tal que Y es la madre de X». Mientras la segunda frase dice que todo el mundo es madre, la primera es la perogrullada de que todo el mundo tiene una madre. No habrá mucha gente que se crea lo que resulta de otras permutaciones posibles de los cuantificadores: «Existe un Y tal que para todo X, Y es la madre de X» («Hay alguien que es la madre de todos») y «Existe un X tal que para todo Y, Y es la madre de X» («Hay una persona tal que todo el mundo es su madre»).

Una situación en la que frecuentemente la gente se hace un lío con los cuantificadores en la vida cotidiana es en las negaciones. Si Jorge le cuenta que todos los habitantes de la isla miden más de 2 metros, y usted quiere negarlo, simplemente dirá que alguien de la isla mide menos de 2 metros; no hace falta que todo el mundo mida menos de 2 metros. Para negar que hay alguien con todos los dientes de oro basta con decir que todo el mundo tiene al menos un diente que no es de oro.

El lenguaje hablado, al contrario que las matemáticas, es a menudo ambiguo, y traducir una frase del lenguaje hablado al lenguaje formal es a menudo un problema espinoso, especialmente si es metafórica. «No es oro todo lo que reluce», por ejemplo, tiene dos formalizaciones completamente distintas, igual que «Se come a todo el mundo» pronunciado por caníbales. Hasta el simple copulativo «es» del español corriente tiene distintas traducciones en la lógica. Compárense sus acepciones en las frases siguientes. En «Paco Rabal es Juncal» tiene el significado de identidad: $R = J$. En «Jorge es tonto» indica predicación: J tiene la propiedad T, o $T(J)$. En «El hombre es racional» significa inclusión: Para todo X, si X tiene la propiedad de ser un hombre, entonces X tiene la propiedad de ser racional, o simbólicamente: $\forall X[H(X) \rightarrow R(X)]$.

Volviendo a la cuestión del principio, notemos que engañar a algunos todo el tiempo es aprovecharse de un grupo especialmente crédulo a los que seguro que se engaña siempre, mientras que engañar a todo el mundo por un rato es perseverar en los propios trucos de timador convencido de que a la larga todo el mundo picará. Por último, la expresión formal de «Todo el mundo es odiado por alguien» es « $\forall Y \exists X H(X, Y)$ ».

E

Más universal aún que la conocida novela *La historia de O* y que las historias kafkianas del Señor K es La historia de E. (Ya sé que es un comienzo estrafalario, pero todo el mundo tiene derecho a sus propias rarezas). Comparable a π en cuanto a importancia matemática y escrito a menudo en una modesta minúscula, e es aproximadamente igual a 2,718 281 828 459 045 235 360 287 471 352 662 497. A continuación esbozaré algunas de sus propiedades.

El número e fue introducido por el matemático suizo Leonhard Euler a mediados del siglo XVIII y, a primera vista, su definición tradicional puede parecer misteriosa. El número se define como el límite de la sucesión de términos $(1 + 1/N)^N$ cuando el entero N se hace más y más grande. Cuando N es 2, la expresión anterior es $(1 + 1/2)^2$, o lo que es lo mismo $(3/2)^2$, es decir 2,25; para N igual a 3, es $(1 + 1/3)^3$, o $(4/3)^3$, que da 2,37; para valores sucesivos de N es $(1 + 1/4)^4$, o $(5/4)^4$, que da 2,44, luego $(6/5)^5$, $(7/6)^6$, ..., $(101/100)^{100}$, etc. El valor de e es el límite de esta sucesión de números. Así pues, es muy próximo a $(10.001/10.000)^{10.000}$, que es igual a 2,718145, pero es aún más próximo a $(1.000.001/1.000.000)^{1.000.000}$.

Aunque sea un tanto abstracta, esta definición encierra la clave del papel de e en los cálculos bancarios y de interés compuesto. (Véase también la entrada sobre *Crecimiento exponencial*). Mil dólares invertidos al 12% se convierten al cabo de un año en $1.000 \times (1 + 0,12)$ dólares. Si se invierten a interés compuesto semestral, se convierten en $1.000 \times (1 + 0,12/2)$ dólares al cabo de seis meses (pues el 12% anual equivale al 6% semestral), y en 1.000

$\times (1 + 0,12/2) \times (1 + 0,12/2)$, o $1000 \times (1 + 0,12/2)^2$ dólares al final del año. Si se trata de interés compuesto trimestral, se convierten en $1000 \times (1 + 0,12/4)$ dólares al final del primer trimestre (pues el 12% anual equivale al 3% trimestral), $1.000 \times (1 + 0,12/4)^2$ dólares al final del segundo trimestre, y en $1000 \times (1 + 0,12/4)^4$ al final del año. Si seguimos así y calculamos el interés compuesto N veces al año, al final de dicho año el dinero se habrá convertido en $1.000 \times (1 + 0,12/N)^N$ dólares. Nótese que, excepto porque tiene 0,12 en vez de 1, el último factor es idéntico a la definición de e . Unos pocos cálculos matemáticos entre bastidores muestran que a interés compuesto diario ($N = 365$) ese dinero se convierte al cabo del año en $1.000 \times e^{0,12}$ dólares y en $1.000 \times e^{0,12T}$ dólares al cabo de T años. A propósito, la función exponencial $Y = e^T$, en términos de la cual se expresa el crecimiento exponencial, es una de las más importantes en matemáticas.

Hay otras definiciones de e , todas ellas equivalentes, por supuesto, y todas ponen de manifiesto (con un poco de abracadabra) el carácter natural de este número. Por ésta y otras razones relacionadas con el cálculo, el número e es la base del sistema de logaritmos *naturales*. Para aclarar esta afirmación hay que extenderse un poco en el análisis de un tema tan desagradable como los logaritmos. El logaritmo *vulgar* o decimal de un número no es más que la potencia a la que hay que elevar 10 para obtener el número en cuestión. El logaritmo decimal de 100 es 2 puesto que $10^2 = 100$ [así pues $\log(100) = 2$]; el logaritmo decimal de 1.000 es 3, pues $10^3 = 1000$, y el logaritmo decimal de 500 es 2,7, pues $10^{2,7} = 500$.

Por su parte, el logaritmo natural de un número es la potencia a la que hay que elevar e para obtener dicho número. Así, el logaritmo natural de 1.000 es aproximadamente 6,9 porque $e^{6,9} = 1.000$ [o de otro modo, $\ln(1.000) = 6,9$]; el logaritmo natural de 100 es 4,6 pues $e^{4,6} = 100$, y el logaritmo natural de 2 es 0,7 porque $e^{0,7} = 2$. Se puede demostrar (una manera matemática de decir «Creedme») que este último número, el logaritmo natural de 2, tiene un papel importante en el mundo de las finanzas: dividiendo 0,7 por el rédito se obtiene el número de años que tarda en doblarse el dinero invertido. Así, con

réditos del 10% ó del 14% ($0,10$ y $0,14$) se tarda respectivamente 7 o 5 años ($0,7/0,1 = 7$ y $0,7/0,14 = 5$). Pero en vez de explicar por qué esto es así, o qué es exactamente lo natural de los logaritmos naturales, explicaré algunos modos en los que el número e se da en otros contextos comunes. (Y, desde luego, como los logaritmos decimales se basan en el hecho accidental de que tengamos 10 dedos, no podemos pretender en modo alguno que sean matemáticamente naturales).

Imaginemos un departamento de una universidad que va a entrevistar sucesivamente a N candidatos para un puesto de profesor ayudante. Al final de cada entrevista, el departamento ha de decidir si el candidato entrevistado es el idóneo. Supongamos que si se descarta a un cierto candidato, no se le puede reconsiderar después, y que, si se llega al último candidato, hay que escogerlo por necesidad. Con el fin de maximizar las probabilidades de escoger al mejor, el departamento decide la siguiente estrategia: escoge cuidadosamente un número $K < N$, entrevista a los primeros K candidatos y los rechaza, y luego sigue con las entrevistas hasta encontrar un candidato mejor que todos los que le han precedido. Y contrata a esa persona.

Esta estrategia no siempre funciona. Unas veces el mejor candidato estará entre las primeras K personas rechazadas, y otras el mejor candidato vendría después del que se ha contratado. Sin embargo, y dadas estas condiciones, se puede demostrar que la estrategia óptima consiste en tomar K igual a $(N \times 1/e)$, donde $1/e$ es aproximadamente $0,37$ o el 37%. Así, si hay 40 candidatos y se entrevistan al azar, la mejor estrategia consiste en rechazar sin más a los 15 primeros (el 37% de 40) y, a partir de ahí, aceptar al primer candidato que sea mejor que todos sus predecesores. La probabilidad de elegir al mejor candidato por este método es también, aunque suene extraño, $1/e$ o el 37%. Ninguna otra estrategia da una probabilidad de éxito mayor del 37%. Argumentos similares se manejan en estrategias parecidas de elección de esposa, aunque en esta situación las condiciones del enunciado del problema son menos naturales.

Tenemos otra aparición inverosímil del número e cuando una secretaria mezcla 50 cartas distintas y sus 50 sobres con las direcciones respectivas. Si mete las cartas en los sobres al azar, se podría preguntar: ¿cuál es la probabilidad de que al menos una carta esté en el sobre que le corresponde?

Por razones no muy fáciles de explicar, el número e interviene también en la solución a este problema. La probabilidad de que haya al menos una coincidencia es $(1 - 1/e)$, o aproximadamente el 63%. Otros enunciados que dan el mismo resultado son el de levantar un par de cartas de dos mazos que se han barajado por separado, o el de los sombreros ordenados al azar y los correspondientes resguardos en el guardarropa de un restaurante.

El número e también aparece inesperadamente en situaciones en las que nos interesamos por el establecimiento de algún récord. A modo de ilustración, imaginemos una región de la Tierra que ha tenido el mismo clima durante eones. Con todo, la pluviosidad anual de esta región presentará fluctuaciones estadísticas. Si tuviéramos que empezar a partir de la pluviosidad del año 1, veríamos que los récords de pluviosidad se dan cada vez menos a medida que pasan los años. La pluviosidad del año 1 constituiría, naturalmente, un récord, y quizá la del año 4 sería superior a la de los tres años anteriores, con lo que se establecería un nuevo récord. Probablemente tendríamos que esperar hasta el año 17 para que la pluviosidad superara la de los 16 años anteriores y se estableciera un nuevo récord. Si siguiéramos registrando las precipitaciones anuales por otros 10.000 años, nos encontraríamos con que sólo se bate el récord de pluviosidad unas 9 veces. Y, si consideráramos un período de un millón de años, probablemente nos encontraríamos con que el récord se bate 14 veces.

No es ninguna coincidencia que la raíz 9.^a de 10.000 y la raíz 14.^a de 1.000.000 sean aproximadamente iguales a e . Si al cabo de N años se ha batido R veces el récord de pluviosidad, la raíz R -ésima de N será una aproximación a e , que será tanto más aproximada cuanto mayor sea N . (El número N ha de ser suficientemente —esto es, enormemente— grande para tener aproximaciones precisas).

A pesar de ser irracional (imposible de expresar como cociente de dos números enteros y, por tanto, no tener una expresión decimal periódica) y trascendente (no es la solución de ninguna ecuación algebraica), e es omnipresente en las fórmulas y teoremas matemáticos. Está íntimamente relacionado con las funciones trigonométricas, las figuras geométricas, las ecuaciones diferenciales, las series infinitas y muchas otras ramas del análisis matemático. La inverosímil trinidad literaria del principio sólo era un intento

egregio de sugerir de una manera no matemática la enorme importancia de e .

Ecuaciones diferenciales

El análisis (el cálculo infinitesimal y sus descendientes) ha sido una de las ramas predominantes de la matemática desde que fue inventado por Newton y Leibniz. Las ecuaciones diferenciales son su núcleo principal. Esta materia ha sido tradicionalmente la clave para comprender las ciencias físicas y, en los temas más profundos sugeridos por ella misma, ha sido el origen de muchos de los conceptos y teorías que constituyen el análisis superior. Es también una de las herramientas prácticas esenciales de que disponen los científicos, los ingenieros, los economistas y otros para manejar tasas de variación. (Una de sus cualidades peor conocidas es el placer menor que produce a algunos estudiantes de segundo año de matemáticas cuando hablan de su curso de *difi-q*.^[5] Una vez conocí a un estudiante de farmacia que escogió esta asignatura porque le encantaba pasarse el día diciendo *difi-q*).

La derivada de una cantidad (véanse las entradas sobre *Cálculo* y *Funciones*) es una función matemática que mide la tasa de variación de dicha cantidad. Por ejemplo, si lanzamos una bola al aire, la derivada de su altura respecto del tiempo es su velocidad. Si C es el coste de producción de X artículos, la derivada de C con respecto a X es el coste marginal de producir el X -ésimo artículo. Las derivadas de orden superior —verbigracia, la derivada de la derivada— miden a qué velocidad cambian las propias tasas de variación. Aunque las ecuaciones en las que intervienen derivadas quizá deberían llamarse ecuaciones derivadas, se llaman ecuaciones diferenciales. Y así como la idea de la resolución de ecuaciones algebraicas es determinar un número a partir de ciertas condiciones que debe satisfacer, resolver una

ecuación diferencial consiste en determinar el valor de una cantidad variable en cualquier instante de tiempo a partir de ciertas condiciones sobre la derivada (y las derivadas de orden superior) de dicha cantidad. Dicho llanamente, el estudio de las ecuaciones diferenciales se ocupa de los métodos y técnicas para determinar el valor de una cantidad en cualquier instante cuando se conoce cómo cambian dicha cantidad y otras relacionadas con ella.

Veamos algunos ejemplos de situaciones que conducen a ecuaciones diferenciales: Nicolae parte al mediodía de Bucarest en dirección oeste y mantiene una velocidad constante de 80 kilómetros por hora; determinar su posición en cualquier instante de esa tarde. Un gran depósito contiene 400 litros de agua en la que se han disuelto 100 kilogramos de sal; si entra agua pura a razón de 12 litros por minuto y la mezcla, que es agitada para que se mantenga uniforme, fluye fuera del depósito a razón de 8 litros por minuto ¿cuánto tiempo ha de transcurrir para que el contenido de sal del depósito sea de 50 kilogramos? Un conejo corre en dirección este a 7 kilómetros por hora, y 1.000 metros al norte de la posición inicial del conejo hay un perro que empieza a perseguirlo a 9 kilómetros por hora dirigiéndose siempre hacia el conejo; determinar la trayectoria seguida por el perro. Una cuerda elástica ideal está atada por ambos extremos y tiramos de ella, encontrar su posición en cualquier instante posterior. [La ecuación correspondiente a este último caso es importantísima en física; para aquellos interesados en ella, es $Y''(X) + KY(X) = 0$, donde $Y''(X)$ (o, en una notación alternativa, d^2Y/dX^2) indica la segunda derivada de $Y(X)$, la desviación de la cuerda relativamente a su posición de equilibrio en cualquier punto X de la misma].

Las derivadas primera y segunda de una cantidad tienen significados marcadamente distintos. Si la cantidad en cuestión es una distancia o una altura, la primera derivada es su velocidad y la segunda su aceleración. Como los problemas de la física no son corrientes en la vida cotidiana, consideraremos el siguiente ejemplo tomado del telediario: Con voz sonora, el comentarista de televisión informa que determinado índice económico sigue subiendo, aunque no tan rápidamente como el mes pasado.

Quizá sin saberlo está diciendo que la derivada del índice es positiva (la tasa de variación del índice es positiva), pero la derivada segunda del índice

es negativa (la tasa de variación de la tasa de variación del índice es negativa). El alza está «llegando a un máximo». Siguiendo bastante más allá por este camino se llega a conjuntos de ecuaciones diferenciales que relacionan los valores, sus tasas de variación y las tasas de las tasas de variación de varios indicadores económicos. Estos conjuntos de ecuaciones constituyen un modelo econométrico y se pueden manejar para hacerse una idea de cómo funciona el mundo real.

Al aplicar las ecuaciones diferenciales a menudo nos interesa introducir más de una variable; a veces no conocemos como cambia una cantidad con respecto al tiempo sino como cambia con respecto a alguna otra variable; muchas situaciones sólo se pueden describir mediante un conjunto de ecuaciones diferenciales interrelacionadas. Los progresos matemáticos realizados en el tratamiento de estos problemas durante los últimos 300 años se cuentan entre las mayores glorias de la civilización occidental. Las leyes del movimiento de Newton, la ecuación del calor de Laplace y la ecuación de ondas, la teoría del electromagnetismo de Maxwell, la ecuación de Navier-Stokes de la dinámica de fluidos y los sistemas depredador-presa de Volterra no son más que una pequeña parte de los muchos frutos que han dado estas técnicas (aunque, tristemente, la mayoría de sus autores son desconocidos para cualquier persona medianamente instruida).

En los últimos tiempos la investigación ha abandonado este campo clásico de las *difi-q*. Ahora se concentra más en las aproximaciones y el cálculo numérico, y no tanto en los métodos tradicionales en los que intervienen los límites y procesos infinitos.

Estadística: dos teoremas

En su libro *Suicide*, el sociólogo francés Emile Durkheim demostró que la incidencia del suicidio en una zona se puede predecir razonablemente sólo en base a los datos demográficos. Análogamente, la tasa de desempleo se puede estimar basándose en muestreos (y otros varios índices económicos). En realidad, muchas predicciones sociológicas y económicas son independientes de las ideas y principios psicológicos y se basan en buena medida en razones probabilísticas. Aunque los sucesos concretos sean difícilmente pronosticables (quién se va a suicidar o quién se quedará sin empleo), los conjuntos grandes de sucesos son en general fáciles de describir de antemano. *Muy* en líneas generales, esto es lo que sugieren dos de los resultados teóricos más importantes de la teoría de la probabilidad y la estadística. (Véanse también las entradas sobre *Media*, *Correlación* y *Probabilidad*).

Concretando un poco más, la ley de los grandes números dice que la diferencia entre la probabilidad de un cierto suceso y la frecuencia relativa con que se produce tiende necesariamente a cero. En el caso de una moneda no cargada, por ejemplo, la ley, descubierta por el matemático suizo Jakob Bernouilli en un trabajo póstumo que fue publicado en 1713, nos dice que se puede demostrar que la diferencia entre $1/2$ y el cociente del número de caras entre el total de lanzamientos se hace arbitrariamente pequeña si aumentamos indefinidamente el número de éstos.

No hay que entender esto como que la diferencia entre los números totales de caras y de cruces irá disminuyendo cada vez más a medida que aumente el número de lanzamientos; normalmente ocurre precisamente todo lo contrario. Si se lanza una moneda 1.000 veces y otra 1.000.000 de veces,

probablemente el cociente del número de caras entre el de lanzamientos sea mucho más próximo a $1/2$ en el segundo caso, a pesar de que la diferencia entre los números de caras y cruces sea también mayor. Las monedas no trucadas se comportan bien en el sentido relativo de los cocientes, pero no en sentido absoluto. Y, contra lo que suponen muchos sabios de salón, la ley de los grandes números no implica la falacia del jugador: que es más fácil que salga cara después de una tira ininterrumpida de cruces. No lo es.

Entre otras creencias justificadas por esta ley tenemos la confianza del experimentador en que la media de un conjunto de medidas de una cierta cantidad se aproximará más al valor real de ésta cuanto mayor sea el número de mediciones. También es la base de la observación razonable de que si se tira un dado N veces, la probabilidad de que la frecuencia con que aparece el 5 sea distinta de $1/6$ disminuye al aumentar N . Al igual que el dado, nosotros, considerados individualmente, tampoco somos predecibles, pero tomados colectivamente sí. La ley de los grandes números sirve de base teórica a la idea intuitiva de que la probabilidad es la guía del mundo. Las clasificaciones de Nielsen en televisión, las encuestas Gallup, las tarifas de seguros y un sinnúmero de estudios sociológicos y económicos ponen de manifiesto una realidad probabilística más confusa que la de las monedas y los dados, pero no menos auténtica.

La otra ley que quiero esbozar aquí se llama teorema central del límite, y dice que la media o la suma de una gran colección de medidas de una magnitud dada cualquiera es descrita por una distribución o curva normal en forma de campana (también llamada a veces curva gaussiana en honor del gran matemático del siglo XIX Karl Friedrich Gauss). Esto ocurre aunque la propia distribución de las medidas individuales no sea normal.

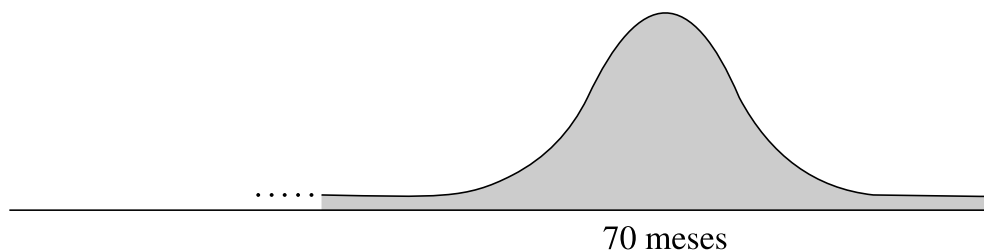
Para ilustrar esto último, imaginemos una fábrica que produce disqueteras para ordenador, y supongamos que el director es un chapucero subversivo que garantiza que aproximadamente el 30% de las disqueteras se estropeen en tan sólo 5 días y que el 70% restante tarden unos 100 meses en fallar. Está claro que la distribución de las vidas de estas disqueteras no obedece a una curva normal, sino a una curva en forma de U con dos picos, uno a los 5 días y otro mayor a los 100 meses.

Supongamos ahora que las disqueteras salen de la línea de montaje en un

orden aleatorio, en cajas de 36. Si nos entretuviéramos en calcular la vida media de las disqueteras de una caja, encontraríamos que es de unos 70 meses, quizá 70,7. ¿Por qué? Si determinamos la vida media de las disqueteras de otra caja de 36, encontraremos de nuevo una vida media de aproximadamente 70 meses, quizá 68,9. De hecho, si examinamos muchas cajas, la media de las medias será muy próxima a 70, y lo que es más importante aún, la distribución de estas medias será aproximadamente normal (en forma de campana), con el porcentaje adecuado de cajas con vidas medias entre 68 y 70, entre 70 y 72, etc.



Distribución en forma de U de las vidas medias de las disqueteras de una caja de 36 típica



Distribución gaussiana normal de las vidas medias de las disqueteras de muchas de esas cajas

Teorema central del límite

El teorema central del límite dice que en una gran mayoría de casos esta situación es la que cabe esperar: que las medias y las sumas de cantidades, que no tienen por qué estar normalmente distribuidas, siguen una distribución normal.

La distribución normal aparece también en el proceso de medida porque las medidas de una magnitud o una característica cualquiera tienden a tener una «curva de error», con forma de campana normal centrada en torno al verdadero valor de dicha magnitud. Otras cantidades que suelen tener una

distribución normal podrían ser las alturas y pesos para una edad específica, el consumo de gas natural de una ciudad en cualquier día dado de invierno, los grosores de piezas, los CI (independientemente de lo que puedan indicar), el número de ingresos en un gran hospital en un día determinado, las distancias de los dardos a la diana, los tamaños de las hojas, de las narices o el número de pasas contenidas en una caja de cereales para el desayuno. Todas estas cantidades se pueden considerar como medias o sumas de muchos factores (genéticos, físicos o sociales) y, por tanto, su distribución normal se basa en el teorema central del límite. Repito, las medias o sumas de una cantidad tienden a estar normalmente distribuidas, aun cuando las cantidades que se promedian (o suman) no lo estén.

Ética y matemáticas

Desde Platón a los filósofos contemporáneos, como John Rawls, pasando por Kant, los moralistas han sostenido la necesidad de unos principios impersonales de la moralidad. La matemática es a veces objeto de mofa por ser una materia impersonal, pero bien entendida, esta impersonalidad es en parte lo que la hace tan útil, incluso en la ética, donde su invocación ha podido parecer rara en principio. Entre otros, el gran filósofo judío-holandés del siglo XVII Spinoza nos dio un ejemplo de esto al escribir su obra clásica *Ética* «al estilo geométrico» de los *Elementos* de Euclides.

Para empezar, un ejemplo deprimente muy alejado de los principios racionalistas y «teoremas» estoicos de Spinoza. Según un informe del UNICEF de 1990, anualmente mueren millones de niños de cosas tan poco graves como sarampión, tétanos, infecciones respiratorias o diarrea. Estas enfermedades se podrían evitar con una vacuna de 1,50 dólares, 1 dólar de antibióticos o 10 centavos de sales hidratantes por vía oral. El UNICEF estima que bastarían 2.500 millones de dólares para salvar las vidas de la mayoría de esos niños y mejorar la salud de muchísimos más. Esta cantidad equivale al presupuesto anual de publicidad de las compañías tabaqueras norteamericanas (cuyos productos, por cierto, matan casi 400.000 norteamericanos cada año, más que los que murieron en toda la segunda guerra mundial), al gasto mensual de los soviéticos en vodka o, más grave aún, a casi el 2% del gasto anual en armamento del propio Tercer Mundo. Además, si se proporcionaran medios para el control de la natalidad a aquellas mujeres que lo desearan —una estimación conservadora da aproximadamente 500 millones— haría disminuir el crecimiento de la

población en un 30%, con lo que la carga financiera que representan los anteriores cambios presupuestarios sería más llevadera. La planificación familiar junto con la garantía de supervivencia de sus hijos comportaría un nuevo recorte de la tasa de crecimiento de la población, pues ésta es mayor en aquellos países con una gran tasa de mortalidad infantil. La aritmética no es complicada, pero resulta esencial para ver la situación desde una buena perspectiva.

La gente siente a menudo aversión a asignar valores numéricos a las vidas humanas o a explicitar determinadas transacciones. Sopesar el coste de la sanidad o el precio del impacto sobre el medio ambiente es siempre una tarea desagradable. A veces, sin embargo, no ser cuantitativo es un modo de falsa piedad que no puede sino oscurecer, y por tanto complicar, las decisiones que nos vemos obligados a tomar. Es ahí donde pueden jugar un papel importante la teoría de la probabilidad y la investigación operacional. Otras veces, me atrevería a añadir, la aritmética económica apropiada es más cantoriana (tanto en el sentido bíblico como en el de la teoría de conjuntos infinitos); esto es, cuando cada vida tiene un valor infinito y es, por tanto, tan valiosa como la suma de cualquier conjunto de vidas también infinitamente valiosas, exactamente igual que $\aleph_0$, el primer número cardinal transfinito de Cantor, que es igual a $\aleph_0 + \aleph_0 + \dots + \aleph_0$. (Véase la entrada sobre *Conjuntos infinitos*).

La necesidad de compromisos no siempre se aprecia en toda su importancia, a pesar de que es algo bastante corriente. En vez de seguir discutiendo sobre ello pondré un par de «aplicaciones» no estándar de la matemática en el campo de la ética. El primero es de naturaleza matemática sólo en un sentido amplio, pero uno de mis propósitos es ampliar la concepción popular de la matemática.

Supongamos que una sociedad ha de tomar una decisión política importante y que decidirse positivamente implica asumir un gran riesgo en el futuro. Si se adopta, esa política implicará inicialmente cierto trastorno — gente que cambia de residencia, mucha inversión en construcción, formación de nuevas organizaciones, ...— pero comportará un aumento del nivel de vida durante 200 o 300 años por lo menos.

En un cierto momento posterior hay, sin embargo, una gran catástrofe, directamente atribuible a la adopción de la política de riesgo, en la que mueren 50 millones de personas. (La decisión podría estar relacionada con el almacenamiento de residuos radiactivos). Ahora bien, como ha indicado el filósofo inglés Derek Parfit, se podría decir que la decisión de adoptar la política de riesgo no fue mala para nadie. La decisión no fue mala, en efecto, para las personas que vieron incrementado su nivel de vida en los siglos anteriores a la catástrofe.

Es más, tampoco fue mala para las personas que murieron en la catástrofe, pues ellos, esos mismos que murieron, no hubieran nacido de no haberse tomado la decisión de seguir la política de riesgo. Ésta, recordémoslo, provocó inicialmente cierto trastorno y la consiguiente alteración del momento en que las parejas existentes concibieron a sus hijos (y a ello deben éstos su ser) y también, debido a que se reunió a diferentes personas que se aparejaron y fueron padres (y de ahí el ser de sus hijos). Con el paso de los siglos, estas diferencias se multiplicaron y es razonable suponer que nadie que vivió el día de la catástrofe habría nacido si no se hubiera adoptado la política de riesgo en cuestión. Las personas que mueran, repitémoslo, deberán su existencia a la toma de esa decisión.

Tenemos pues un ejemplo de decisión, tomar el camino del riesgo, que parece ser claramente mala —conduce a la muerte de 50 millones de personas— y, sin embargo, no es mala (discutiblemente) para nadie. Lo que nos hace falta es algún o algunos principios morales a cuya luz se pueda rechazar la política de riesgo.

Un candidato a esta categoría es el principio utilitarista del filósofo del siglo XIX Jeremy Bentham: «El mayor bien para el mayor número». Sin embargo, el principio sólo es esquemático y una interpretación precisa del mismo, que permitiera crear una especie de «cálculo moral» es algo que, gracias a Dios, los filósofos morales no han logrado todavía. Kant y otros pensaban que cualquier principio moral debería ser universal, algo que fuera útil en abstracto, pero que no sirviera demasiado al pasar a los detalles. En vez de entretenerme en la inmensa literatura relativa a éstos y otros enfoques de la ética, introduciré aquí un fragmento cuasimatemático que resultará ilustrativo independientemente de cuál sea nuestro enfoque de los principios

éticos.

Formulado originalmente en términos de presos, el dilema del preso es de una generalidad difícilmente sobreestimable. Supongamos que dos hombres sospechosos de un delito importante son detenidos mientras cometen una falta menor. Son separados e interrogados, y a cada uno se le da la posibilidad de confesar el delito importante, implicando con ello a su cómplice, o permanecer callado. Si ambos permanecen en silencio, les caerá un año de prisión a cada uno. Si uno confiesa y el otro no, el que confiesa será recompensado con la libertad, mientras que al otro le caerá una condena de cinco años. Si ambos confiesan, pueden esperar que les caigan tres años de cárcel. La opción cooperativa es permanecer callado, mientras que la individualista es confesar.

El dilema consiste en saber qué es lo mejor para ambos en conjunto. Dado que permanecer callados y pasar un año en prisión deja a cada uno de ellos en manos de la peor de las posibilidades, pecar de incauto y que le caiga una condena de cinco años, lo más probable es que ambos confiesen y que cada uno pase tres años en la cárcel.

La gracia del dilema no es, naturalmente, el interés que podamos tener en mejorar el sistema jurídico-penal, sino el hecho de presentar el mismo esqueleto lógico que muchas situaciones que hemos de afrontar en nuestra vida cotidiana. (Véase la entrada sobre *Teoría de juegos*). Tanto si somos ejecutivos en un mercado competitivo, cónyuges en un matrimonio o superpotencias en una carrera armamentista, nuestras opciones pueden formularse en términos similares a los del dilema del preso. Aunque no siempre haya una respuesta correcta, generalmente las partes implicadas salen mejor paradas en conjunto si cada una resiste la tentación de traicionar a la otra y coopera con ella o le es leal. Si ambas partes persiguen exclusivamente su propio interés, el resultado es peor para ambas que si cooperan. La mano invisible de Adam Smith, encargada de que la búsqueda del provecho particular comporte el bienestar colectivo, está, al menos en estas situaciones, totalmente artrítica.

Las dos partes del dilema del preso se pueden generalizar a circunstancias en las que participa mucha gente, donde cada individuo tiene la opción de hacer una contribución minúscula al bien común u otra mucho mayor en

beneficio propio. Este dilema del preso a muchas bandas es útil para modelizar situaciones en las que está en juego el valor económico de «intangibles» como el agua limpia, el aire o el espacio. Como en buena medida casi todas las transacciones sociales tienen en sí algún elemento del dilema del preso, el carácter de una sociedad queda reflejado en cuáles son las transacciones que llevan a la cooperación entre las partes y cuáles no. Si los miembros de una «sociedad» concreta nunca se comportan cooperativamente, sus vidas serán probablemente, en palabras del filósofo inglés Thomas Hobbes, «solitarias, pobres, repugnantes, brutales y cortas».

¿Es a consecuencia de alguna teoría moral que uno elige la opción cooperativa en situaciones del tipo del dilema del preso? Por lo que yo sé, no. De hecho, se puede hacer una sólida defensa de la opción individualista, al menos algunas veces. No es irracional ni inmoral defenderse a uno mismo. No hay ninguna teoría ética establecida que obligue ni prohíba adoptar la opción cooperativa, y lo mismo ocurre con muchas otras acciones, y esto me lleva a mi última observación.

Cualquier teoría moral, convenientemente formalizada y cuantificada, está sujeta a las limitaciones impuestas por el primer teorema de incompletitud de Gödel (véase la entrada sobre *Gödel*), según el cual todo sistema formal que sea lo suficientemente complejo, ha de contener forzosamente enunciados de los que no se puede probar ni la verdad ni la falsedad. De acuerdo con esto, tenemos una base teórica para la observación racional de que siempre habrá actos que ni nos están prohibidos ni estamos obligados a ellos por nuestros principios, independientemente de cuáles sean éstos o de que estén reforzados por nuestros propios temores, valores y compromisos idiosincrásicos. Esto podría tomarse como un argumento matemático en pro de la necesidad de la «ética de la situación» y demuestra la insuficiencia de una aproximación a la ética exclusivamente axiomática.

Fermat y su último teorema

Muchas personas que se entretienen jugando con números han descubierto que los números 3, 4 y 5 tienen la interesante propiedad de que $3^2 + 4^2 = 5^2$. Algunos sin duda han descubierto también que hay otras ternas de números con esta misma propiedad. Otros dos ejemplos son las ternas 5, 12, 13 y 8, 15, 17, que satisfacen las ecuaciones $5^2 + 12^2 = 13^2$ y $8^2 + 15^2 = 17^2$. Se ha demostrado que hay una infinidad de ternas pitagóricas como éstas.

Al tratarse de una propiedad tan simple y natural, los matemáticos se han preguntado si sería posible generalizarla. Han investigado, en particular, si hay ternas de números enteros X , Y , y Z tales que satisfagan que $X^3 + Y^3 = Z^3$, y no han encontrado ninguna. Han buscado ternas de números X , Y , y Z que fueran solución de $X^4 + Y^4 = Z^4$, y tampoco han encontrado ninguna. De hecho ni los matemáticos ni nadie que se distraiga jugando con números ha encontrado nunca un conjunto de tres números enteros X , Y y Z ni ningún número entero N mayor que 2 tales que satisfagan la ecuación $X^N + Y^N = Z^N$.

El gran matemático del siglo XVII Pierre Fermat (que debería ser recordado por sus importantes contribuciones a la teoría de los números, la geometría analítica y la teoría de la probabilidad) escribió que este fracaso tenía sus buenas razones: no existen números enteros X , Y y Z tales que $X^N + Y^N = Z^N$ para ningún número entero N mayor que 2. Fermat lo demostró para $N = 3$ y en una página de un texto clásico griego sobre teoría de los números anotó que tenía una elegante demostración del teorema general, pero

que en el margen de la página no tenía suficiente espacio para reproducirla. Probablemente su demostración fuera incorrecta, pero nadie más la encontró, ni nadie ha conseguido tampoco, en los tres siglos transcurridos desde su muerte, encontrar otra demostración del resultado general que se conoce como el último teorema de Fermat. Se han encontrado muchos resultados parciales (el dominio de valores de N para los que no existen soluciones ha ido ampliándose sin cesar) y recientemente se han propuesto algunas argumentaciones que casi demuestran el teorema.^[6] La importancia matemática del teorema no es inherente al teorema en sí (que sólo es una curiosidad), sino que radica en toda la teoría algebraica de los números descubierta e inventada a raíz de los esfuerzos por demostrarlo. El teorema se parece más a una mota de carbonilla en el ojo que a un diamante auténtico o, cambiando la metáfora del carbono por otra de silicio, un grano de arena en una ostra cuya irritante presencia acabará por producir una perla. (Una cuestión teórica más sustanciosa fue planteada a principios de este siglo por el matemático alemán David Hilbert como parte de su famosa lista de problemas pendientes. Además de las interrogantes sobre la hipótesis del continuo de Cantor, la consistencia de la aritmética y otras cuestiones abstractas, Hilbert se preguntó si existía algún método general para determinar si los polinomios arbitrarios de varias variables —como $3X^2 + 5Y^3 - 21X^5Y = 12$ — tenían soluciones enteras. Recientes avances de la lógica han demostrado que puede que dicho método no exista).

Aunque todavía no haya sido demostrado, hay consenso en que el último teorema de Fermat es cierto.^[*] Con todo, si resulta ser falso, lo único que hace falta para demostrar su falsedad es un sólo contraejemplo: una terna de números enteros X , Y y Z y algún número entero N mayor que 2 tal que: $X^N + Y^N = Z^N$. ¿Alguna sugerencia para probar? [Y, como despedida, unos cuantos ramalazos de lógica cuya omisión pudiera ser especialmente provechosa para el lector. El párrafo anterior demuestra que, debido a su forma, si el último teorema de Fermat es falso, entonces es refutable. Por tanto, si no es refutable, es verdadero. Así pues, si es indecidible (ni demostrable ni refutable), es verdadero. Y podremos concluir que, *si* el último teorema de Fermat es una proposición aritmética indecidible (y hay

muchas: véase la entrada sobre *Gödel*), entonces es verdadero. Este razonamiento hipotético no arroja ninguna luz sobre la verdad o falsedad del teorema. Es curioso, sin embargo, que a partir de un conocimiento metamatemático de una insuficiencia de la aritmética, podamos deducir la veracidad de un simple enunciado aritmético].

La filosofía de la matemática

¿Qué son los números, los puntos y las probabilidades? ¿De qué naturaleza es la verdad matemática? ¿Por qué es útil la matemática? Éstas son algunas de las preguntas cuyas respuestas no encontrará aquí. No obstante, intentaré presentar un par de asuntos relacionados con ellas.

El observador más despreocupado se da cuenta de que los teoremas matemáticos no se confirman del mismo modo que las leyes físicas. Parecen ser verdades necesarias, mientras que las afirmaciones de las ciencias empíricas (la física, la psicología y la cocina) parecen depender completamente de la manera en que el mundo es realmente. Al menos desde un punto de vista conceptual, las leyes de los gases de Boyle y la historia del Imperio austrohúngaro podrían haber sido fácilmente otras, mientras que no se puede decir lo mismo de la afirmación $2^5 = 32$.

Pero ¿de dónde viene la certeza y la necesidad de la verdad matemática? Los matemáticos en activo no suelen preocuparse por este asunto pero, si se les aprieta, la mayoría contestarán probablemente algo así como: los objetos matemáticos existen independientemente de nosotros y las afirmaciones sobre ellos son verdaderas o no independientemente de nuestro conocimiento y de nuestra capacidad de demostrarlas. Imagino que tales objetos existen en algún mundo platónico más allá del tiempo y el espacio. Pero si es así, ¿cómo encontramos verdades sobre tales objetos y los hechos que a ellos atañen?

La respuesta de Immanuel Kant era que la matemática (o al menos sus axiomas fundamentales) era conocible *a priori* por la sola intuición y que su necesidad era evidente. Los intuicionistas contemporáneos, sin suscribir las

ideas kantianas acerca del espacio, el tiempo y el número, también basan la necesidad de la matemática en la indudabilidad de las actividades mentales simples. Algunos incluso llegan a desacreditar las demostraciones de la existencia de un objeto a menos que se dé un procedimiento constructivo que permita encontrarlo.

A muchos otros filósofos de la matemática les molesta tanto el subjetivismo de Kant como la insostenibilidad del platonismo ingenuo. Los llamados logicistas, Bertrand Russell, Alfred North Whitehead y Gottlob Frege, trataron de demostrar que la matemática se podía reducir a la lógica y por lo tanto era tan cierta como la simple proposición «A o no A», y que en el fondo los enunciados matemáticos no eran sino maneras tortuosas de decir «A o no A». Su esperanza era encontrar así una garantía de la certeza de los enunciados matemáticos, pero no acabaron de lograrlo del todo. Lo que llamaban lógica contenía ideas de la teoría de conjuntos que eran precisamente tan problemáticas como los enunciados matemáticos que se seguían de ellas.

La respuesta convencionalista a la pregunta fue que la matemática alcanzaba su necesidad por convenio, por *fiat* y por definición. Sus verdades no eran más que cuestión de convenio y, por tanto, no eran ni más ni menos oscuras que el hecho de que 5 pesetas sean un duro. Con un enfoque parecido, los filósofos formalistas sostenían que los enunciados matemáticos no hacen referencia a nada, sino que sólo son sucesiones de símbolos gobernadas por reglas, exactamente igual que las reglas del ajedrez rigen el movimiento de las piezas sobre el tablero. Que el movimiento del caballo sea dos cuadros en una dirección y uno en perpendicular es algo necesario pero no misterioso.

La suficiencia con que dan por concluida la cuestión estas últimas posiciones es atractiva al principio, pero son absolutamente incapaces de explicar lo que el Premio Nobel Eugen Wigner dio en llamar «la irrazonable eficacia de la matemática» en la descripción de la realidad. Otros filósofos replican que la correspondencia entre las estructuras matemáticas y la realidad física no es en absoluto «irrazonable». Es, según ellos, no muy distinta de la razonable correspondencia entre los distintos sentidos biológicos (vista, oído, olfato, gusto y tacto) y los aspectos de la realidad

física. La percepción matemática sería una especie de sexto sentido abstracto.

De dónde viene la necesidad de la matemática y si los números son construcciones mentales, facetas de una realidad idealizada o sólo símbolos que se rigen por unas reglas, son temas que han resonado bajo diversas formas a lo largo y ancho de toda la historia de la filosofía. En la Edad Media, por ejemplo, los protagonistas de la batalla eran los idealistas, los realistas y los nominalistas, y sus posiciones acerca de la naturaleza de los universales como la Rojez y la Triangularidad eran en cierto modo análogas a las que sostienen los intuicionistas, logicistas (o platónicos) y formalistas de hoy. (Véanse también las entradas sobre *Geometría no euclídea*, *Probabilidad* y *Sustituibilidad*).

Estos asuntos trascienden la matemática. Están íntimamente relacionados, por ejemplo, con la distinción filosófica entre verdades analíticas y verdades sintéticas. Una proposición analítica es verdadera en virtud del significado de las palabras que la forman, mientras que la veracidad de una proposición sintética se da en virtud del modo como son las cosas. (Un tipo especial de afirmaciones analíticas, las lógicamente válidas, son verdaderas en virtud del significado de las partículas lógicas «y», «o», «no», «si..., entonces...», «algún» y «todo». Las afirmaciones que son verdaderas en virtud de las cuatro primeras de estas partículas lógicas se llaman tautologías. Véanse las entradas sobre *Cuantificadores* y *Tautologías*).

Así, «Si Pedro huele mal y es desdentado, entonces huele mal» es una proposición analítica, mientras que «Si Pedro huele mal, entonces es desdentado» es sintética. Otros ejemplos del mismo antagonismo son «Los solteros son hombres no casados» frente a «Los solteros son hombres lujuriosos» y «Los ovnis son objetos voladores que no han sido identificados» frente a «En los ovnis viajan unas criaturas verdes». Los filósofos cuentan las verdades matemáticas entre las analíticas y la mayoría de las demás entre las sintéticas, y, aunque no sea inamovible, esta distinción es un recurso conceptual práctico. Cuando el pomposo médico de Molière dice que la poción soporífera es eficaz gracias a su poder adormecedor, enuncia una frase vacía y analítica, y no de tipo objetivo y sintético.

Al moverse en tomo a cuestiones de verdad trascendente y certidumbre, la filosofía de la matemática tiene también una cierta resonancia con el

pensamiento religioso. ¿Por qué razón si no, un agnóstico convencido como yo habría empleado con tanta frecuencia palabras como «divino», «sacerdotal», «cielo», «pureza» y «reverencia» en estas entradas? Las semejanzas son generalmente metafóricas, pero a veces las metáforas determinan las actitudes y los actos.

Por último, sea cual sea el «ismo» al que uno se adhiera (o del que se proteja), los teoremas de incompletitud de Gödel (véase la entrada sobre Gödel) barajaron los naipes filosóficos y obligaron a todas las partes a recomponer sus respectivos juegos. La existencia de proposiciones indemostrables indica, por ejemplo, que la verdad de las mismas no puede radicar únicamente en su demostración a partir de los axiomas. Más aún, la misma consistencia de una teoría matemática es una de las proposiciones que no se pueden demostrar sino que simplemente hay que suponer (o aceptar por la fe, si se prefiere este lenguaje).

Folklore matemático

Aunque en principio, la frase «folklore matemático» pueda sonar un tanto extraña, algo así como «cuentos de hadas por ordenador» o «parábolas electrónicas», la matemática, al igual que otras disciplinas, tiene sus propios cuentos e historias, a menudo apócrifas, que constituyen un marco de referencia común y sirven para definir un carácter específico.

La presentación de muchos de los teoremas, ejemplos y principios que se han discutido aquí ha ido acompañada de una pequeña narración (el misticismo de Pitágoras, los puentes de Euler, el último teorema de Fermat, la paradoja de Russell), y en esta entrada sólo quiero esbozar algunas historias arquetípicas más. Las palabras «esbozar» y «algunas» están vigiladas aquí en un sentido literal. No pretendo ser exhaustivo (lo cual a veces sólo es un eufemismo para dejar exhausto al lector).

Empezaré con dos famosos episodios de Arquímedes, el mayor matemático, físico e inventor de la Antigüedad. Se cuenta que mientras se remojaba en una bañera descubrió el principio que lleva su nombre: todo cuerpo sumergido en un fluido experimenta un empuje hacia arriba igual al peso del fluido que desaloja. Con la alegría del descubrimiento, salió corriendo desnudo a la calle gritando, «¡Eureka! ¡Eureka!» (en español «¡Lo encontré! ¡Lo encontré!»). La segunda anécdota sobre Arquímedes también realza su devoción por el saber. Absorto en un diagrama matemático dibujado en la arena, le pasó totalmente desapercibido que un soldado romano detrás de él le había ordenado que no continuara y, por ello, el soldado le mató. Este acto ilustra, tan bien como cualquier otro, la relación entre las culturas griega y romana en la Antigüedad.

De gustos matemáticos bastante más refinados que Arquímedes, el

matemático inglés del siglo XX G. H. Hardy se ufanaba muchísimo de la inutilidad de la teoría de los números. Es muy conocida la siguiente conversación entre él y su protegido, el matemático indio Srinivasa Ramanujan. Estando de visita en la habitación de este último en el hospital, Hardy dijo que 1729, la matrícula del taxi que le había traído, era un número bastante soso. A lo que Ramanujan replicó casi inmediatamente: «¡No, Hardy! ¡No, Hardy! Es un número muy interesante. Es el menor número que se puede expresar como suma de dos cubos de dos maneras distintas». ($9^3 + 10^3 = 1^3 + 12^3 = 1.729$).

El tema de la ineptitud, o al menos la falta de interés por las cuestiones mundanas, está también presente en muchas de las anécdotas que se refieren al fundador de la cibernética, el matemático del MIT^[7] Norbert Wiener. Se cuenta que tenía una vista y/o una memoria tan malas que se le asignó un estudiante de doctorado para asegurar que llegara a sus destinos. Otra anécdota sobre Wiener ilustra la tendencia al elitismo que parece caracterizar a muchos matemáticos. Estaba en cierta ocasión dando la clase a un curso de doctorado y enseguida descubrió que sólo un estudiante de la primera fila seguía todos los detalles de su exposición. A partir de entonces, Wiener dio la clase hablando directamente a este estudiante. Pero un buen día el estudiante faltó y, al no verlo en primera fila, Wiener se fue inmediatamente refunfuñando que no había nadie en clase.

El potencial intimidatorio de la matemática y los matemáticos lo ilustra una conversación entre el prolífico matemático suizo del siglo XVIII, Leonhard Euler, y el filósofo y enciclopedista francés, Denis Diderot. Antes de una discusión teológica de la que no habría salido muy bien, Euler pidió a Diderot que contestara a una fórmula matemática irrefutable y sin embargo irrelevante: «Señor, $(a + b)^n/n = X$; por tanto, Dios existe. ¡Responda!». Diderot, estupefacto, no supo replicar. (Véase la entrada sobre *QED*).

Raramente se ve a los matemáticos como románticos. Con todo, en 1832 el brillante algebrista francés de veintiún años, Evariste Galois, murió en un duelo por una prostituta. Se dice que (como siempre la historia tiene varias versiones) Alfred Nobel, el inventor de la dinamita y fundador del premio que lleva su nombre, había estipulado que no hubiera un premio Nobel de

matemáticas para desquitarse del amante de su mujer, Mittag-Leffler, que probablemente lo habría ganado en la época inicial de los premios.

No hay una gran abundancia de enemistades personales rencorosas entre matemáticos, y las que hay suelen tener una componente matemática importante. Por ejemplo, los rencores entre los matemáticos alemanes del siglo XIX, Leopold Kronecker y Georg Cantor, fundador éste último de la teoría de conjuntos, giraban en buena parte alrededor de las distintas concepciones del infinito por parte de ambos. Kronecker tenía una idea pitagórica finitista y decía, «Dios hizo los enteros y el resto es obra del hombre». Cantor, por su parte, trató con toda clase de conjuntos y construcciones trascendentes. Los ataques de Kronecker al brillante pero hipersensible Cantor quizá contribuyeran a las depresiones de éste y su posterior internamiento en una institución mental. Existen antagonismos similares, aunque más benignos, entre los matemáticos puros y los aplicados, entre los algebristas y los analistas, y entre los lógicos y los demás matemáticos.

Más típico es el folklore matemático consistente en anécdotas de figuras carismáticas recogidas por los propios matemáticos. Luego se cuentan y vuelven a contar como si fueran viejos chistes (y son una de las pocas buenas razones para asistir a los congresos de matemáticos). Se cuenta, por ejemplo, que Kurt Gödel se resistió durante varios años a hacerse ciudadano norteamericano porque había hallado una contradicción lógica en la Constitución.

Otro ejemplo típico se refiere a John von Neumann, que era considerado por algunos la persona más inteligente que haya existido jamás. Se le planteó el problema de un pájaro que vuela entre dos trenes que se acercan. El pájaro vuela a 150 kilómetros por hora y los trenes, que inicialmente están a 540 kilómetros de distancia, van uno al encuentro del otro a 80 y 40 kilómetros por hora, respectivamente. La pregunta es ¿qué distancia recorre el pájaro antes de morir aplastado entre los dos trenes? La manera laboriosa de resolver el problema consiste en calcular las longitudes de los sucesivos vuelos entre los trenes antes de que éstos choquen y luego sumar la serie resultante. La manera fácil es darse cuenta de que los trenes chocan al cabo de 4,5 horas ($540 \text{ kilómetros} / 120 \text{ kilómetros por hora}$) y, por tanto, el recorrido total del

pájaro es $4,5 \text{ horas} \times 150 \text{ kilómetros por hora} = 675 \text{ kilómetros}$. Cuando Von Neumann soltó 675 casi inmediatamente, quien había planteado el problema se echó a reír y comentó que sin duda Von Neumann sabía el truco, a lo que, según se dice, éste replicó: «¿Qué truco? ¿Hay algo más fácil que sumar la serie?».

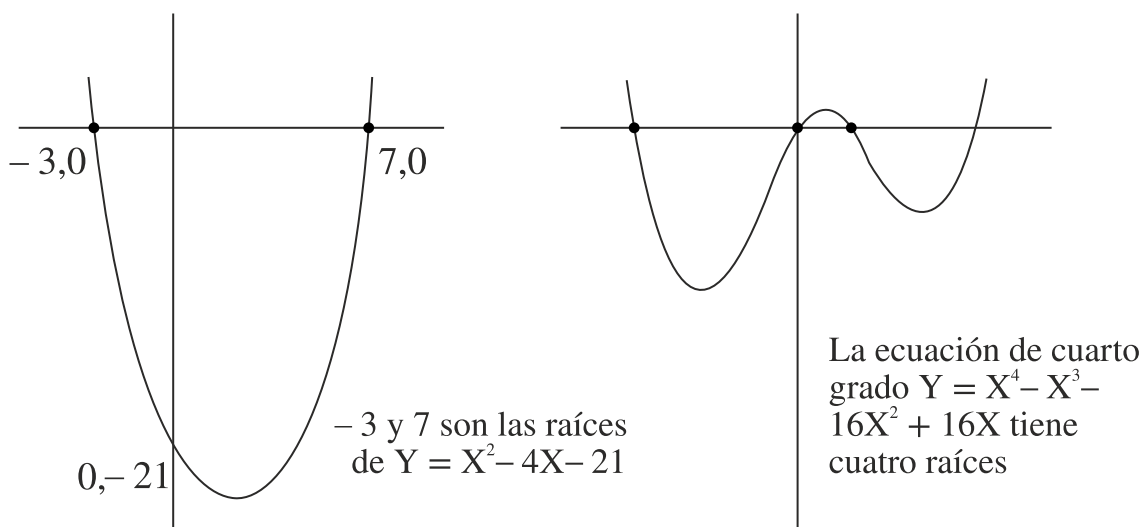
Muchas historias contemporáneas parecen perpetuar la idea de que los matemáticos son una raza aparte, ya sea de imbéciles pies planos (como el matemático proverbial que dice A, escribe B y quiere decir C cuando la verdadera conclusión es D) o de calculadores brillantes y oscurantistas irrelevantes. En una cita clásica, Arquímedes decía que si le daban un punto de apoyo, una palanca lo suficientemente larga y un lugar donde situarse, podría mover el mundo. La cita sugiere la naturaleza teórica, la fuerza práctica y las aspiraciones trascendentales de la matemática y los matemáticos. La idea matemática expresada en ella, el concepto de proporcionalidad, es fundamental, como lo es también el concepto de recurrencia, que es la idea que se expresa en la siguiente historia. Compárese, no obstante, el antiguo ideal con el más moderno estereotipo siguiente.

Un psicólogo pregunta a un ingeniero qué haría si se declarara un pequeño fuego y hubiera un jarro de agua sobre la mesa. El ingeniero responde obedientemente que apagaría las llamas con el agua. A renglón seguido el psicólogo se vuelve hacia un matemático y le pregunta qué haría si se iniciara un pequeño fuego y hubiera un jarro de agua en el alféizar de la ventana. El matemático contesta que trasladando el jarro del alféizar a la mesa se reducía el problema al anterior, cuya solución ya se conocía.

La fórmula de la ecuación de segundo grado y otras fórmulas

La fórmula de la solución de la ecuación de segundo grado es uno de los primeros teoremas de álgebra que se demuestran en el instituto. Permite resolver fácilmente ciertas ecuaciones y parece definir bastante bien lo que muchas personas creen que es la matemática. Para ellos todo el tinglado consiste simplemente en sustituir números en esas fórmulas o quizás en factorizar polinomios. Aunque desde luego no es así, parece que hay personas a las que el hecho de disponer de una fórmula en la que sustituir datos o un polinomio que factorizar les otorga una cierta sensación de aptitud.

Las ecuaciones de segundo grado son ecuaciones, como $X^2 - 4X - 21 = 0$ o $3X^2 + 7X - 2 = 0$, en las que la variable aparece elevada al cuadrado. Son ecuaciones que aparecen a menudo en física, ingeniería y en otras disciplinas. Si, por ejemplo, desde un tejado de 80 metros de altura se tira una piedra hacia arriba con una velocidad de 20 metros por segundo, podría interesarnos saber que la piedra caerá al suelo al cabo de T segundos, siendo T una solución de $-5T^2 + 20T + 80 = 0$ (el 5 es debido al efecto de la gravedad).



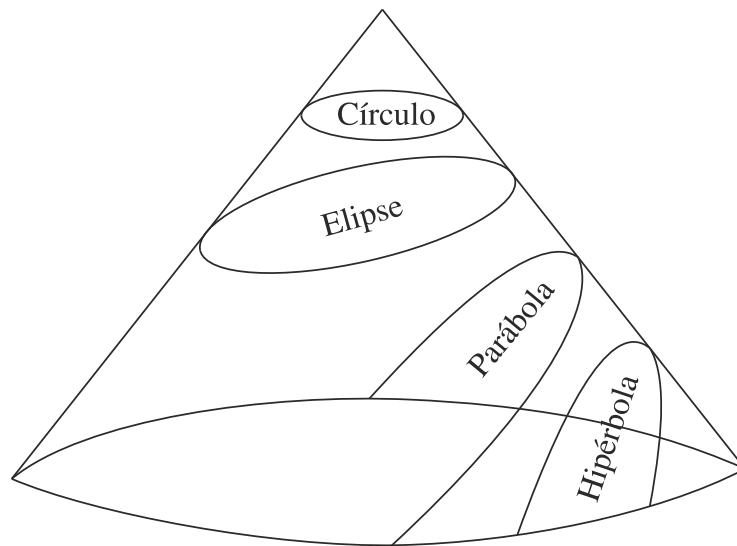
Representación gráfica de ecuaciones

Lo que se necesita en tales situaciones es encontrar las raíces de estas ecuaciones, los números que sustituidos en lugar de X (o T) hacen verdaderas las igualdades expresadas por dichas ecuaciones. Las raíces de la primera de las ecuaciones escritas más arriba son -3 y 7 ; las de la segunda son aproximadamente $-2,59$ y $0,26$; y las de la tercera, $-2,47$ y $6,47$. Hay varias técnicas para encontrar estas raíces (entre las que se cuentan el tanteo, la factorización y tirar piedras desde los tejados), pero ninguna es en general más eficaz que la famosa fórmula de la ecuación de segundo grado: $X = \frac{-B \pm \sqrt{B^2 - 4AC}}{2A}$. (Véase la entrada sobre *Números imaginarios y números negativos*). La fórmula nos da las dos raíces de $AX^2 + BX + C = 0$ en función de los números A , B y C , que en las tres ecuaciones anteriores son respectivamente iguales a 1 , -4 y -21 , a 3 , 7 y -2 , y a -5 , 20 y 80 . Si se sustituyen estos números en la fórmula, se puede ver (una vez pasada la bruma del esfuerzo de cálculo) que las raíces de cada una de las ecuaciones son iguales a los valores dados más arriba.

Con frecuencia, además de estos números que se obtienen al igualar a 0 una expresión cuadrática, nos interesan también los valores de dicha expresión para valores de X arbitrarios. Así, si el coste C de producir X artículos viene dado por $1.000 + 20X + 0,3X^2$ (1.000 pesetas de coste fijo, 20 pesetas por cada artículo más un término cuadrático, $0,3X^2$, que empieza

siendo pequeño pero que aumenta rápidamente con X y que refleja los costes de almacenamiento y oficina que comportan los grandes inventarios), tendremos que utilizar la función cuadrática $C = 1.000 + 20X + 0,3X^2$ y su gráfica. O si nos interesa la relación entre $X^2 - 4X - 21$ y un número arbitrario Y , tendremos que estudiar la función $Y = X^2 - 4X - 21$ y su gráfica. Un rasgo típico de estas funciones cuadráticas es que sus gráficas tienen forma parabólica. En el segundo caso la gráfica corta al eje de las X donde $Y = 0$, en -3 y 7 .

Si examinamos la gráfica de una ecuación cuadrática general de dos variables, $AX^2 + BXY + CY^2 + DX + EY + F = 0$, encontraremos la versión moderna de las antiguas secciones cónicas de Apolonio. Para distintos valores de los números A , B , C , D , E y F tenemos ecuaciones cuyas gráficas son círculos, elipses, parábolas e hipérbolas. Éstas son precisamente las figuras que se forman por intersección de un cono y un plano. Según sea la inclinación del plano con respecto al eje del cono obtendremos una u otra sección cónica.



Las secciones cónicas

Las gráficas de las funciones de grado superior, como la cúbica (donde la variable está elevada al cubo: $Y = 4X^3 - 5X^2 + 13X - 7$) y la cuártica (en la que aparecen cuartas potencias de la variable: $Y = 6X^4 + 5X - 9$), no son tan

naturales desde el punto de vista geométrico. Pero podemos encontrar las raíces de estas ecuaciones, como en el caso de las de segundo grado, mediante unas fórmulas, más complicadas, descubiertas por Gerónimo Cardano, Niccolo Tartaglia y otros matemáticos italianos del siglo XVI. Parecería natural, pues, suponer que todas las ecuaciones de grados superiores (de quinto, sexto y hasta de enésimo grado) se pueden resolver del mismo modo.

El matemático francés del siglo XIX Evariste Galois demostró, sin embargo, que eso no era cierto. No todas las ecuaciones algebraicas pueden resolverse sustituyendo datos en una fórmula consistente en sumas, restas, multiplicaciones, divisiones, potencias y raíces. Para las ecuaciones de grado superior al cuarto no existen tales fórmulas. Galois atacó el problema por un camino abstracto que no pasaba por el estudio concreto de las raíces de la ecuación (o función polinómica), sino que se basaba en la consideración de unas estructuras consistentes en las distintas permutaciones de las raíces. (Una analogía que no sería demasiado engañosa podría ser el estudio de las familias y sus crisis basado en la investigación de las relaciones entre sus miembros y su dinámica, en vez del examen detallado de cada uno de ellos). Para los matemáticos de la época esas ideas resultaban incomprensibles, pero desde hace ya bastante tiempo las mismas forman parte de los fundamentos del álgebra abstracta y son uno de los lazos entre esta materia y la que se estudia en el instituto.

Fractales

I magine que se encuentra al pie de una montaña sin vegetación. Que la ha subido y bajado y que, según sus estimaciones, la distancia andada es aproximadamente 10 kilómetros. Ahora bien, ¿qué pasaría si un gigante de 60 metros de altura hubiera de recorrer el mismo camino de ida y vuelta a la cumbre? Quizás sólo tendría que andar 5 kilómetros. Al ser tan alto podría saltar directamente algunos obstáculos que usted había tenido que subir y bajar. O, por el contrario, piense en un insecto que se arrastrara por el mismo camino. Quizá recorrería 15 kilómetros, pues tendría que subir y bajar por las piedras y pequeños obstáculos que usted simplemente había saltado.

Análogamente, suponga que un pequeño animal del tamaño de una ameba hubiera de culebrear por el mismo camino de ida y vuelta. Quizá recorrería 20 kilómetros, pues tendría que subir y bajar por grietas y protuberancias de las rocas y los guijarros, tan pequeñas que hasta el insecto habría podido saltar. Así pues, llegamos a la conclusión un tanto extraña de que la distancia hasta la cima de la montaña depende en buena medida de quien la recorre. Y otro tanto ocurre con el área de la ladera de la montaña: el animalito del tamaño de una ameba lo encontraría un dominio más espacioso para pasearse que el gigante, que con sus zancadas obvia los pequeños detalles de la superficie. Cuanto mayor es el escalador, más corta es la distancia y menor el área de la ladera. Ésta es una característica de los fractales, y la ladera de una montaña es un buen ejemplo de ellos.

El tronco de un árbol, por poner otro ejemplo típico de fractal, se ramifica en un número característico de ramas, cada una de las cuales, a su vez, se ramifica en el mismo número de ramas menores, que a su vez se dividen en el mismo número de ramas aún menores, y así hasta llegar al nivel de las

ramitas. ¿Qué tiene esto en común con la superficie de una montaña?

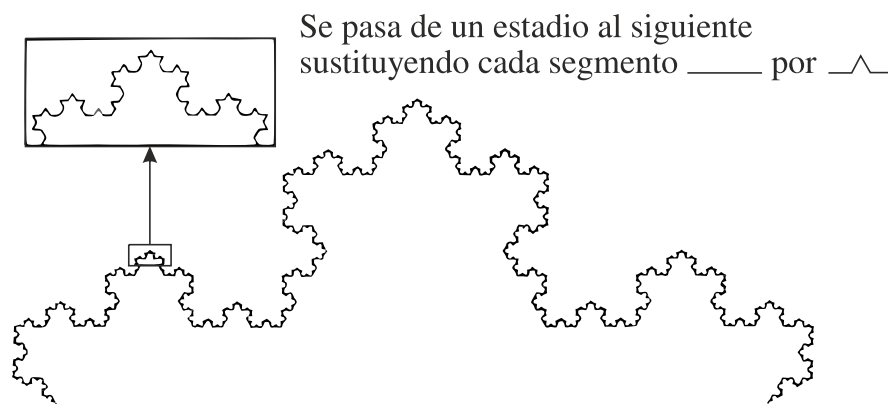
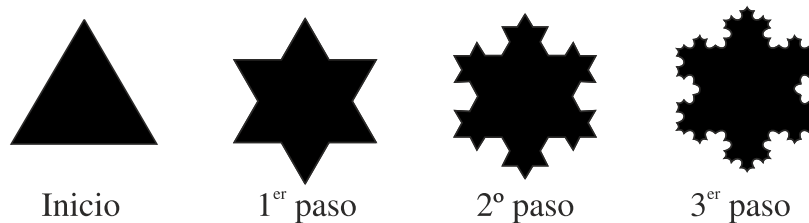
Antes de pasar a la definición, consideremos una línea costera, un ejemplo más debido al matemático Benoît Mandelbrot, descubridor de la geometría fractal. Una estimación de la longitud de la Costa Este de Estados Unidos desde un satélite, por ejemplo, podría dar unos 4.500 kilómetros, más o menos. Si, por el contrario, nos basáramos en el estudio de mapas detallados que muestren los muchos cabos y golfos que se encuentran a lo largo de la costa, quizá nuestra estimación de su longitud aumentaría hasta los 13.500 kilómetros. Si no tuviéramos nada que hacer en un año y decidiéramos andar desde Maine a Miami manteniéndonos siempre a una distancia máxima de un metro o dos del Atlántico, la distancia que recorreríamos se aproximaría más a los 27.000 kilómetros. Aparte de los cabos y golfos, habríamos de tener en cuenta también los entrantes y salientes que no figuran en los mapas normales. Por último, si pudiéramos convencer a un insecto para que recorriera la costa (quizá nuestro amigo escalador prefiera estar al nivel del mar), dándole instrucciones para que no se separara del agua más de la distancia de un guijarro, quizás encontraríamos que la longitud de la costa es de casi 45.000 kilómetros. La costa es un fractal.



Imágenes cada vez más detalladas de la Costa Este de Estados Unidos

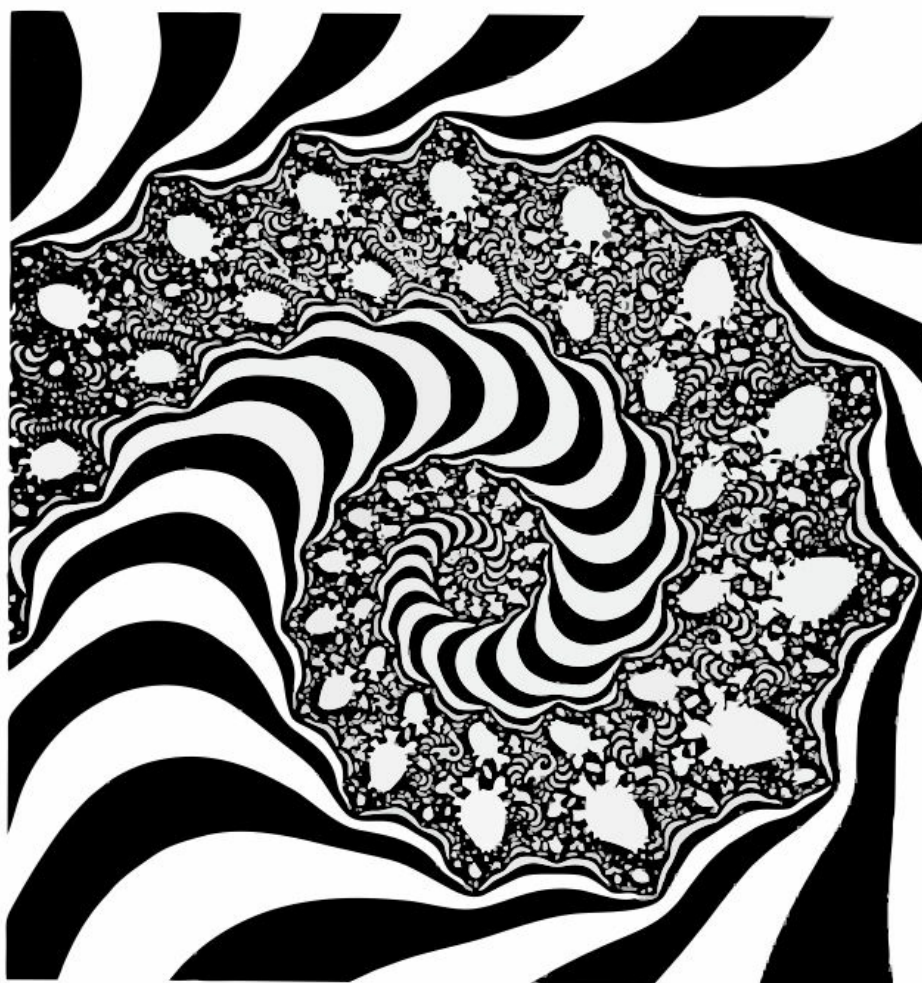
Y también lo es una famosa curva descubierta en 1906 por el matemático sueco Helge von Koch. Koch empezó con un triángulo equilátero y substituyó

cada lado por un segmento con una protuberancia en forma de triángulo equilátero en su tercio central. Repitió este procedimiento una y otra vez y en el límite obtuvo una extraña curva infinitamente enroscada con forma de copo de nieve.



Ampliación de un estadio avanzado de la curva copo de nieve de Koch

¿Qué es un fractal? Es una curva o una superficie (o también un sólido o un objeto de más dimensiones) que presenta una complejidad mayor, aunque parecida, a medida que lo contemplamos de más cerca. La costa, por ejemplo, tiene una forma dentada característica a cualquier escala que la miremos; esto es, tanto si usamos imágenes de satélite para esbozar toda la costa como si empleamos la información mucho más detallada de una persona que recorra andando una pequeña sección de la misma. La superficie de la montaña presenta aproximadamente el mismo aspecto, tanto si se la mira desde los 60 metros del gigante como desde la altura del insecto. La ramificación del árbol tiene la misma apariencia para nosotros que para los pájaros o incluso para los gusanos y los hongos en el caso límite ideal de ramificación infinita. Igual que ocurre con la curva de Koch.



Ampliación de parte de un fractal debido a Benoît Mandelbrot

Por otra parte, como ha señalado también Mandelbrot, las nubes no son circulares ni elípticas, la corteza de un árbol no es suave, el relámpago no sigue una línea recta, y los copos de nieve no son hexagonales (aunque tampoco se parecen a las curvas de Koch). Antes bien, estas formas y muchas otras que se dan en la naturaleza son casi fractales y presentan unos entrantes, salientes, abolladuras, mellas y zigzags característicos a casi cualquier escala. A más ampliación se observan unas convoluciones semejantes y cada vez más complejas. Incluso hay una manera natural de asignar una dimensión fractal a esas figuras. Los fractales que se emplean en los modelos de costas tienen dimensiones comprendidas entre 1 y 2 (más que una recta y menos que un plano), mientras que los que sirven para hacer modelos del relieve tienen

dimensiones entre 2 y 3 (más que un plano y menos que un sólido). Las fotografías de la NASA indican que la dimensión fractal de la superficie terrestre es 2,1, mientras que la de Marte, con una topografía más rugosa y más «lanuda», es 2,4. El término «fractal», acuñado por Mandelbrot en 1975, es una expresión apropiada para esas figuras *fragmentadas*, *fracturadas* y autosemejantes de dimensión *fraccionaria*.

Además de ser de gran ayuda en infografía, donde se usan para producir paisajes y formas naturales verosímiles, los fractales surgen frecuentemente al analizar una estructura fina: de la superficie de los electrodos de una pila, del interior esponjoso de los intestinos y del tejido del pulmón, de la variación temporal de los precios de una determinada mercancía o en la difusión de un líquido a través de una arcilla semiporosa. Con su bella e intrincada complejidad a todos los niveles y escalas de ampliación, los fractales juegan un papel cada vez más importante en la teoría del caos (véase la entrada sobre *La teoría del caos*), donde se emplean para describir el conjunto de posibles trayectorias de un sistema dinámico. Su grotesca elegancia es también patente en contextos puramente matemáticos. Una ecuación dada, por ejemplo, divide el plano en varias regiones, según la raíz a la que converja el método de Newton aplicado a cada uno de sus puntos. Las fronteras que separan estas regiones son fractales asombrosamente *complejos*.

Puede que algún día los novelistas descubran que un análogo de los fractales en el «espacio psíquico» puede servir para captar la estructura fracturada aunque coherente de la consciencia humana, cuyo foco de interés puede pasar casi instantáneamente de las trivialidades de un instante a las verdades eternas del siguiente, conservándose la misma persona en los distintos niveles. (Véase la entrada sobre *La conciencia humana y su naturaleza fractal*). En este sentido, las transcripciones literales de las conversaciones ordinarias son muy reveladoras. Las paradas, los arranques, las elipsis, la rara sintaxis, las referencias vagas, las digresiones sin motivo y los repentinos cambios de dirección no tienen nada en común con la aséptica versión «lineal» que normalmente sale de la imprenta. Quizás haya algún modo de que los conceptos anteriores sean útiles también en psicología cognitiva. La dificultad de una disciplina, por ejemplo, se podría tomar como un fractal, de modo que los más brillantes y/o capaces pudieran superar con

grandes zancadas cognitivas las pequeñas dificultades que otros, menos dotados, han de escalar pacientemente.

Funciones

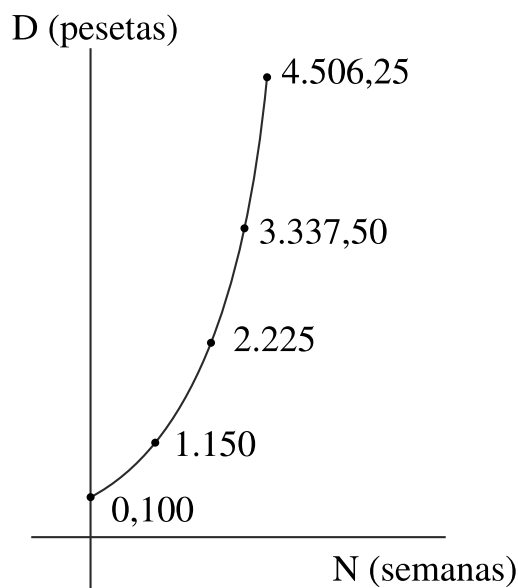
El concepto de función es muy importante en matemáticas, pues representa de una manera formal la idea de poner en correspondencia una cantidad con otra. El mundo está lleno de cosas que dependen de, son función de o están asociadas a otras cosas (de hecho, se podría argumentar que el mundo *consiste* sólo en tales relaciones), y nos enfrentamos al problema de establecer una notación útil para esta dependencia matemática. Los siguientes ejemplos sirven para ilustrar una notación corriente. Las gráficas y las tablas nos proporcionan otras maneras de indicar estas relaciones. (Véase la entrada sobre *Geometría analítica*).

Consideremos un pequeño taller que se dedica a fabricar sillas. Sus costes son 80.000 ptas. (para gastos de equipo, pongamos por caso) y 3.000 ptas. por silla fabricada. Así, la relación entre el coste total, T , y el número de sillas fabricadas, X , viene dada por la fórmula $T = 3.000 X + 80.000$. Si queremos recalcar que T depende de X , decimos que T es función de X y denotamos simbólicamente esta asociación por $T = f(X)$. Si se fabrican 10 sillas, el coste es 110.000 ptas.; si se fabrican 22, el coste sube hasta 146.000 ptas. La función f es la regla que asocia 110.000 a 10 y 146.000 a 22, lo cual se indica escribiendo $f(10) = 110.000$ y $f(22) = 146.000$. ¿Cuánto es $f(37)$?

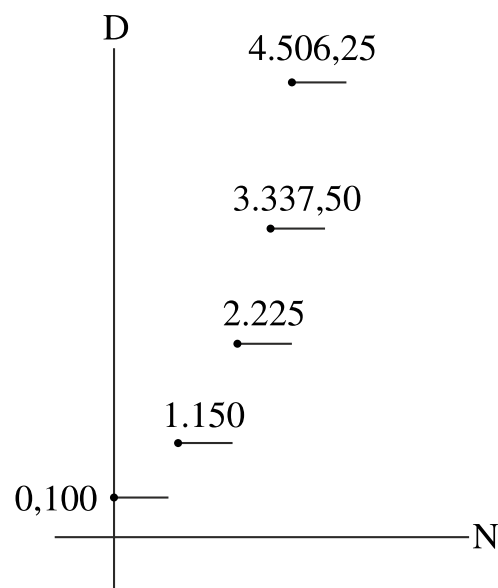
La temperatura Celsius C se puede obtener a partir de la temperatura Fahrenheit F restando 32 a ésta y multiplicando la diferencia por $5/9$. En forma de ecuación tenemos $C = 5/9 (F - 32)$. Así, unos fríos 41° Fahrenheit se convierten en unos igualmente fríos 5° Celsius, mientras que unos suaves 86° Fahrenheit se traducen en otros igualmente suaves 30° Centígrados. Si sustituimos la temperatura Fahrenheit en esta fórmula podemos encontrar siempre la temperatura Celsius correspondiente. Como antes, si lo que

queremos es recalcar que C depende de F , diremos que C es función de F y denotaremos esta relación por $C = h(F)$. (Las gráficas de esta función y la anterior son líneas rectas). La función h es la regla que asocia 5 a 41 y 30 a 86, y esta correspondencia se expresa simbólicamente escribiendo $h(41) = 5$ y $h(86) = 30$. ¿Cuánto es $h(59)$?

O imagine que usted es un usurero que presta 100 ptas. a alguien y le dice que la cantidad que le adeuda aumentará en un 50% cada semana. Revisando las cuentas con sus socios, usted entiende que la cantidad, D , que le debe su amigo al cabo de N semanas es igual a $100 \times (1,5)^N$; esto es, $D = 100(1,5)^N$. Está claro que D es función de N , cosa que indicamos por $D = g(N)$ (o mediante la gráfica de la función, una curva que crece exponencialmente). Está claro que $g(1) = 150$, $g(2) = 225$ y $g(3) = 337,50$. (Si usted es benévolo y sólo añade los intereses a intervalos semanales, la gráfica consistirá en una sucesión de escalones crecientes exponencialmente).



Gráfica de la cantidad debida, D , como función exponencial del tiempo de préstamo, N ; $D = 100(1,5)^N$



Gráfica de la cantidad debida si entre dos incrementos semanales los intereses permanecen constantes

O considere el siguiente ejemplo extraído de la física. Desde un tejado de

80 metros de altura sobre el suelo, se lanza una bola verticalmente hacia arriba con una velocidad inicial de 20 metros por segundo. Confíe en la palabra de Newton y acepte que la altura A de la bola sobre el nivel del suelo viene dada por la fórmula $A = -5T^2 + 20T + 80$, donde T es el número de segundos transcurridos desde el instante en que se lanzó la bola. Como la altura depende del tiempo, A es función de T y se escribe $A = s(T)$. Si sustituimos $T = 0$ en la fórmula, confirmamos que en el instante inicial $A = 80$. Dos segundos más tarde, $T = 2$, encontramos por sustitución en la misma fórmula que $A = 100$. Por tanto $s(0) = 80$ y $s(2) = 100$. ¿Cuánto es $s(5)$? ¿Por qué es menor que $s(2)$?

Las funciones h , g y s anteriores son funciones lineal, exponencial y cuadrática, respectivamente, mientras que $p(X) = 3\text{tg}(2X)$ y $r(X) = 7X^5 - 4X^3 + 2X^2 + 11$ se llaman, respectivamente, trigonométrica y polinómica. Aunque las funciones no siempre están definidas por fórmulas y ecuaciones, ni tienen por qué indicar necesariamente relaciones entre números. Por ejemplo, si $m(\text{Elena}) = \text{rojo}$, $m(\text{Rebeca}) = \text{amarillo}$, $m(\text{Marta}) = \text{moreno}$, $m(\text{Jorge}) = \text{negro}$, $m(\text{Dorita}) = \text{dorado}$ y $m(\text{Pedro})$ no está definido, no es difícil adivinar que m es la regla que a cada persona le asigna el color de su cabello y que Pedro es calvo. Así pues, $m(X)$ denota simplemente el color de cabello de X . Análogamente, $p(X)$ se podría definir como el autor de X y $q(X)$ podría ser la capital de estado más próxima a X . En tal caso, $p(\text{Guerra y paz}) = \text{Tolstoi}$ y $q(\text{Filadelfia}) = \text{Trenton, N. J.}$

En los ejemplos expuestos, el número de sillas fabricadas, la temperatura Fahrenheit, el número de semanas hasta que se salda la deuda, el número de segundos transcurridos desde que se lanza la bola y el nombre de la persona son lo que se llama la variable independiente. El coste total, la temperatura Celsius, la cantidad adeudada, la altura de la bola y el color del cabello de la persona son lo que se llama la variable dependiente. Una vez se ha fijado el valor de la variable independiente, el de la variable dependiente queda totalmente determinado y se dice que ésta es función de aquélla.

Cuando tenemos cantidades que dependen de más de una cantidad —esto es, cuando tratamos con funciones de más de una variable— se usan variantes de la misma notación. Si por ejemplo $Z = X^2 + Y^2$, entonces cuando $X = 2$ e

$Y = 3$ tenemos $Z = 13$, y si queremos resaltar la dependencia de Z con respecto a X e Y , escribimos $Z = f(X, Y)$ y $13 = f(2, 3)$.

La notación de la dependencia funcional es como una contabilidad, pero una contabilidad imprescindible. Nos permite expresar relaciones en forma abreviada. Gracias a ella podemos disponer fácilmente de una buena parte de la flexibilidad y potencia del análisis matemático.

[Respuestas a las preguntas: $f(37) = 191.000$; $h(59) = 15$; $s(5) = 55$ y $s(2) = 100$; en $T = 2$ la bola está subiendo y en $T = 5$, bajando].

Geometría analítica

Como sucede con muchos descubrimientos fundamentales, la idea que llevó a la geometría analítica es, vista retrospectivamente, simple y evidente para todo el mundo, y en especial para los taxistas. En el lenguaje de los taxistas, esta idea se puede expresar así: a cada cruce le corresponde una calle y una avenida, y a cada pareja formada por una calle y una avenida le corresponde un cruce. En una formulación más matemática se dice que a cada punto le corresponde un par ordenado de números y a cada par ordenado de números le corresponde un punto.

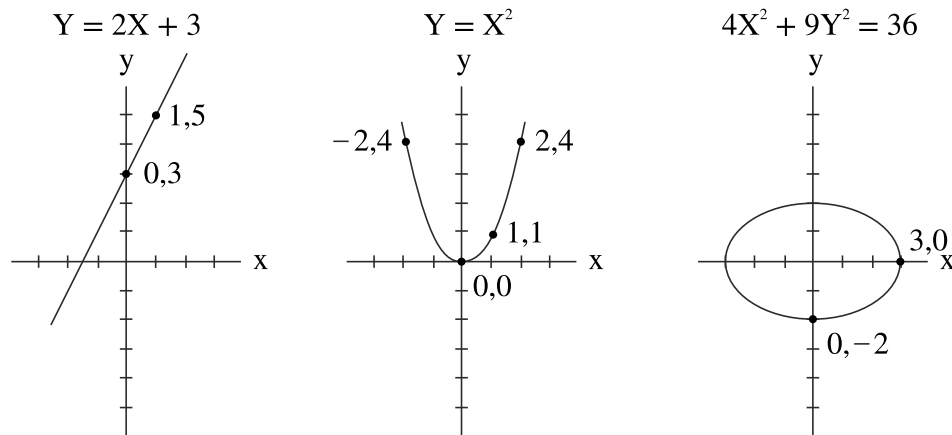
Inventada independientemente por el filósofo y matemático francés René Descartes y su compatriota, el abogado Pierre Fermat, a principios del siglo XVII, la geometría analítica relaciona el álgebra y la geometría por medio de las correspondencias anteriores. El punto $(3, 8)$, por ejemplo, es el punto que se encuentra 3 unidades a la derecha y 8 unidades por encima de un punto fijo que se llama origen. Los números 3 y 8 se llaman respectivamente coordenada X y coordenada Y del punto, e indican un punto distinto del designado por los números $(8, 3)$, cuya posición está 8 unidades a la derecha del origen y 3 unidades por encima. Las coordenadas del origen son, como parece bastante natural, $(0, 0)$. Para el taxista es muy clara la importancia del orden de los números, pues no es lo mismo $(3, 8)$, Tercera Avenida y Calle 8, que $(8, 3)$, Octava Avenida y Calle 3. La intersección de la primera avenida en dirección norte-sur (o eje Y en términos matemáticos) con la primera calle en dirección este-oeste (o eje X) se toma como punto de referencia del taxista u origen.

Aunque se trata de puntos completamente distintos, tanto $(3, 8)$ como $(8, 3)$ están a la derecha (al este) y por encima (al norte) del origen. Por

convenio, los puntos que están a la izquierda (al oeste) o por debajo (al sur) del origen se representan por coordenadas negativas. Así, un punto 5 unidades a la izquierda y 11 unidades por encima del origen tiene coordenadas $(-5, 11)$, otro 5 unidades a la izquierda y 11 unidades por debajo tiene coordenadas $(-5, -11)$ y uno que esté 5 unidades a la derecha y 11 unidades por debajo, $(5, -11)$. (No se ha de tomar Nueva York como modelo pues allí la numeración de las avenidas aumenta según se va hacia el oeste).

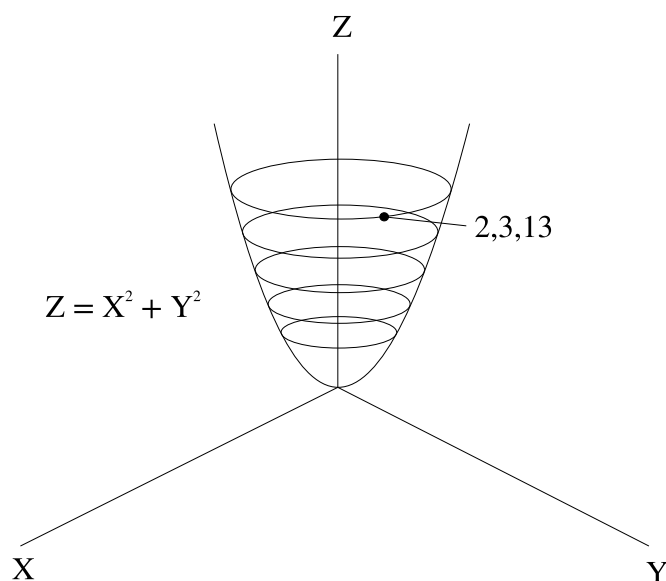
Bueno, ¿y qué? Si la geometría analítica sólo consistiera en esto, las universidades aprobarían automáticamente la asignatura a los taxistas que se matricularan en ellas. Pues bien, después de asociar puntos a pares ordenados de números, Descartes y Fermat observaron además, y esto es crucial, que las ecuaciones algebraicas se correspondían con figuras geométricas. Por poner un ejemplo sencillo, notamos que el conjunto de puntos cuyas coordenadas X e Y satisfacen la ecuación $Y = X$ forman una línea recta. Esto es, los puntos $(1, 1)$, $(2, 2)$, $(3, 3)$, etc. (que satisfacen la ecuación $Y = X$) están sobre una recta que forma un ángulo de 45° con el eje X . Análogamente, los cruces de la Calle 1 y la Primera Avenida, la Calle 2 y la Segunda Avenida, la Calle 3 y la Tercera Avenida, etc., están sobre una recta (o paseo diagonal) que forma un ángulo de 45° con la primera calle este-oeste de referencia.

La gráfica de la ecuación $Y = 2X$ se determina de un modo parecido, y es una recta que forma un ángulo mayor con el eje X . Como los puntos $(1, 2)$, $(2, 4)$, $(3, 6)$ satisfacen la ecuación (cada uno cumple que si se sustituye X por la coordenada X e Y por la coordenada Y , resulta una igualdad), las intersecciones de la Primera Avenida y la Calle 2, la Segunda Avenida y la Calle 4, la Tercera Avenida y la Calle 6, etc., caen sobre esa línea. Por supuesto, no hace falta que nos limitemos a los números enteros; $(1,8, 3,6)$ también cae sobre dicha línea, igual que $(2,7, 5,4)$, aunque ninguno de estos números corresponde a una intersección de calles y avenidas. Análogamente, la ecuación $Y = 2X + 3$ tiene una gráfica que pasa por los puntos $(0, 3)$, $(1, 5)$, $(2, 7)$, $(3, 9)$, etc. Representando gráficamente dos ecuaciones con el mismo par de ejes coordenados, podemos encontrar donde se cortan. Este punto es lo que se llama solución simultánea de las dos ecuaciones, si bien, como la matemática es atemporal, sería más apropiado hablar de solución «homolocal».



Recta, parábola y elipse

Ecuaciones más complicadas dan lugar a curvas más interesantes (cuyos puntos siempre se pueden determinar encontrando pares ordenados de números que satisfagan la ecuación en cuestión). La gráfica de la ecuación $Y = X^2$ es una curva denominada parábola, la de $X^2 + Y^2 = 9$ es una circunferencia, y la de $4X^2 + 9Y^2 = 36$, una elipse. Por medio de estas técnicas que desarrollan esta aproximación algebraico-geométrica, los problemas de geometría se pueden reformular en términos algebraicos y, a la vez, se puede dar una interpretación geométrica a las relaciones algebraicas. La unificación resultante sirvió de marco para el desarrollo posterior del cálculo, en el siglo XVII, y dio lugar a una especie de *lingua franca* vigente aún hoy en día. (Vigente aún, podría añadir, incluso entre los taxistas de Nueva York, ya hablen persa, español, griego, hebreo o ruso. El origen de este símil de taxistas se remonta a una experiencia que viví cuando conducía un taxi siendo estudiante graduado en la Universidad de Wisconsin, en Madison. El encargado insistía en emplear la hora militar y las coordenadas matemáticas para guiar hacia sus destinos a los conductores neófitos, muchos de los cuales eran forasteros. El problema era que Madison es una ciudad emparedada entre cuatro lagos y no hay ningún sistema de coordenadas rectangulares, sino solamente diagonales curvilíneas cortando entre este lago y aquella península. Los conductores solían perderse hasta que llegó otro encargado, desconocedor de la geometría analítica, que les dirigía de un modo más convencional, mediante semáforos, almacenes y gasolineras).



Paraboloide

La generalización a tres dimensiones es inmediata. Así como cualquier punto del plano se puede pensar como un par ordenado de números, cualquier punto del espacio se puede pensar como una terna ordenada de números. El punto (4, 7, 5), por ejemplo, está 4 unidades al este, 7 unidades al norte y 5 unidades por encima de un cierto punto de referencia que tiene coordenadas (0, 0, 0); las coordenadas indican simplemente las distancias sobre los tres ejes X, Y y Z, respectivamente (en vez de sólo dos como en el plano). Volviendo a nuestra perspectiva de taxista, podemos imaginar que (4, 7, 5) está en el quinto piso de un edificio que hay en el cruce de la Cuarta Avenida con la Calle 7, mientras que (4, 7, -1) está en el sótano de este mismo edificio. Los puntos que satisfacen ecuaciones de tres variables describen superficies en el espacio en lugar de curvas en el plano. La gráfica de $Z = X^2 + Y^2$, por ejemplo, es un paraboloide, que tiene una forma parecida a una taza de café redondeada y sin asa, mientras que la gráfica de $X^2 + Y^2 + Z^2 = 25$ es una esfera.

Hay generalizaciones para espacios de más dimensiones, y la idea de asignar números a cada una de las varias dimensiones de una entidad es algo archirrepetido en muchos contextos ajenos a la matemática. Así, el «vector» (4, 7, -1, 14) podría indicar el ya citado sótano del edificio de la Cuarta

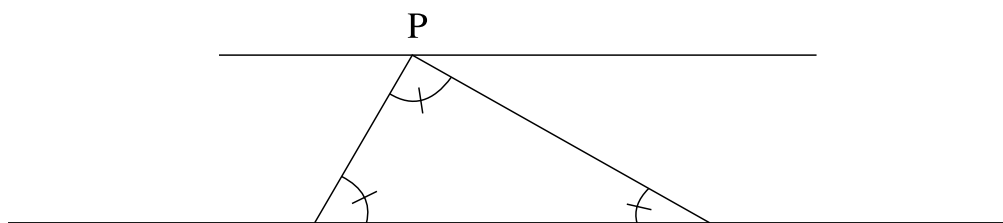
Avenida y la Calle 7 a las 2 de la tarde (14 horas), mientras que (4, 10, 9, 16) indicaría el 9.º piso de un edificio 3 manzanas más al norte y 2 horas después.

Al igual que nuestro sistema de numeración indo-arábigo, la geometría analítica y sus retoños son algo aparentemente tan natural, y por ello tan aceptado de antemano, que hay que hacer un auténtico esfuerzo para recordar que son inventos humanos y no aspectos innatos de nuestra naturaleza conceptual o biológica. (El gran filósofo alemán Immanuel Kant quizá no estaría de acuerdo con esta última observación, pero como todavía no hemos llegado a la geometría no euclídea, no vamos a prestarle atención por el momento).

Geometría no euclídea

En la confusión de propiedades relativas a los triángulos y paralelogramos, semejanza y congruencia, áreas y perímetros, se olvida a veces el carácter deductivo de la geometría. Muchas de esas propiedades fueron descubiertas por los egipcios y los antiguos babilonios a partir de rutinas prácticas de la agrimensura, el comercio y la arquitectura. Los griegos demostraron que todas podían deducirse a partir de unas pocas de ellas. Es fácil formular la idea fundamental. Se escogen algunas suposiciones geométricas «evidentes» que se llaman axiomas y a partir de ellas se deducen, con la única ayuda de la lógica, toda una serie de enunciados geométricos que se llaman teoremas. En su tratamiento del tema, Euclides escogió cinco axiomas (en realidad son diez, pero sólo cinco de ellos eran geométricos) y dedujo el bello y prestigioso cuerpo de teoremas que se conoce como geometría euclídea. (Véase la entrada acerca del *Teorema de Pitágoras*).

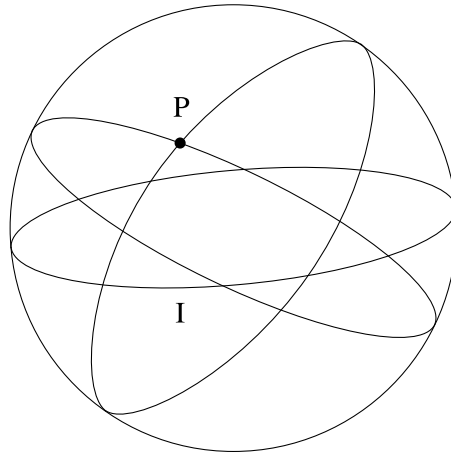
Uno de los cinco axiomas de Euclides es el conocido como postulado de las paralelas. Decía (y sigue diciendo) que por un punto exterior a una recta dada se puede trazar una recta paralela a la dada, y sólo una. Una conocida consecuencia del postulado de las paralelas es el teorema que dice que la suma de los ángulos interiores de un triángulo es siempre 180° .



La suma de los ángulos de un triángulo es 180° . Por P pasa una recta paralela a I, y una sola

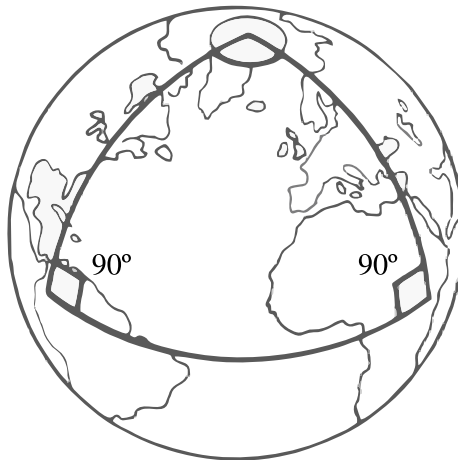
Como no parecía tan intuitivo como los otros cuatro axiomas, a lo largo de la historia los matemáticos han tratado esporádicamente de deducirlo a partir de ellos. Se inventaron todos los métodos que uno pueda imaginar pero nunca dieron con una demostración. Este fracaso, unido a la naturalidad de los otros axiomas, parecía conferir a la geometría euclídea un cierto carácter absoluto. A lo largo de un par de milenios reinó como un monumento al sentido común y la verdad eterna. Immanuel Kant llegó a decir incluso que la gente sólo podía pensar el espacio en términos euclídeos. Por fin, en el siglo XIX los matemáticos János Bolyai, Nikolái Lobachevski y Karl Friedrich Gauss se dieron cuenta de que el quinto postulado de Euclides es independiente de los otros cuatro y no se puede deducir de ellos. Es más, comprendieron que se puede sustituir el quinto postulado por su contrario y tener también un sistema geométrico consistente.

Para ver esto, obsérvese que es perfectamente posible interpretar los términos fundamentales de la geometría de un modo completamente distinto sin salirse, no obstante, de los límites de la lógica más estricta. Exactamente igual que «todos los A son B y C es un A, luego C es un B» nos sirve de justificación para los argumentos más dispares, según sean las interpretaciones de A, B y C, también los términos de la geometría euclídea se pueden interpretar de un modo nada convencional sin dejar de llevarnos a teoremas válidos. Por ejemplo, en vez de las interpretaciones habituales podemos llamar «plano» a la superficie de una esfera, «punto» a un punto sobre una esfera y «línea recta» a los círculos máximos de la esfera (cualquier arco de circunferencia alrededor de la esfera que la divide en dos mitades). Si adoptamos estos significados, la «línea recta» sigue siendo el camino más corto entre dos puntos (sobre la superficie de la esfera) y tenemos una interpretación de la geometría que cumple todos los axiomas de Euclides excepto el postulado de las paralelas. También se cumplen todos los teoremas deducibles a partir de estos cuatro axiomas.



No hay ninguna «línea recta» paralela a la «línea recta» I que pase por P.
En esta interpretación se cumplen todos los demás axiomas de Euclides.

Comprobando los axiomas, observamos que por dos «puntos» cualesquiera pasa una «línea recta», pues dados dos puntos cualesquiera sobre la superficie de una esfera hay un círculo máximo que pasa por ellos. (Nótese que el círculo máximo que pasa por Los Angeles y Jerusalén cruza por Groenlandia y es la distancia más corta entre ambas ciudades). Dado un «punto» y una distancia cualesquiera, hay un «círculo» sobre la superficie de la esfera con centro en ese punto y cuyo radio es esa distancia (que no es más que un círculo normal sobre la superficie de la esfera). Y también, dos «ángulos rectos» cualesquiera son iguales, y cualquier «segmento de recta» (un pedazo de círculo máximo) se puede prolongar indefinidamente.

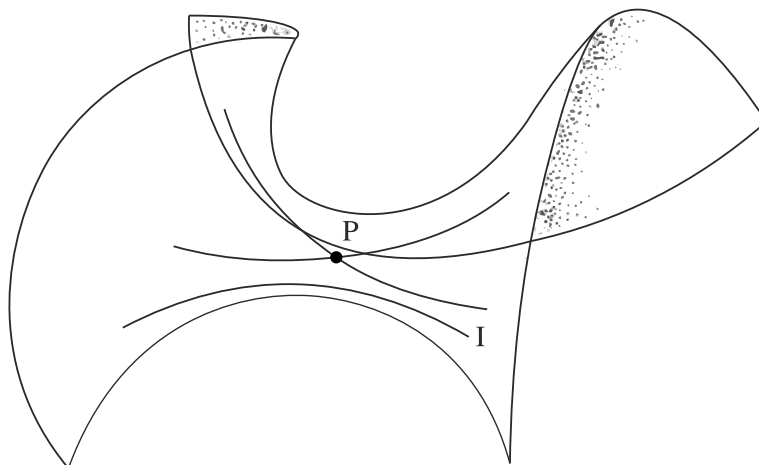


Los «segmentos de recta» que unen Kenia, Ecuador y el Polo Norte

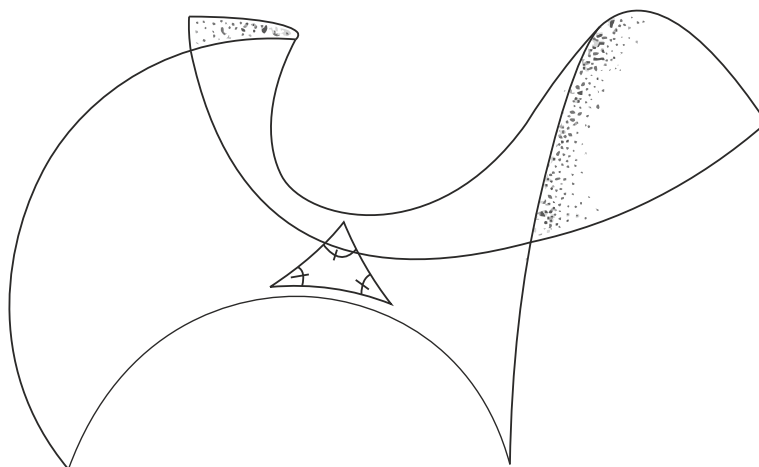
forman un «triángulo» cuyos ángulos suman más de 180°

Sin embargo, el axioma de las paralelas no es válido en esta interpretación particular de los términos, pues dada una «línea recta» y un punto exterior a ella, no hay ninguna «línea recta» paralela a la dada que pase por dicho punto. A modo de ejemplo, tomemos el ecuador como la «línea recta» y la Casa Blanca en Washington como «punto» exterior a la misma. Cualquier «línea recta» que pase por la Casa Blanca será un círculo máximo que divida la Tierra en dos mitades y, por tanto, cortará necesariamente el ecuador, con lo que no podrá serle paralela. Otra anomalía de esta interpretación es que la suma de los ángulos de un triángulo es siempre mayor que 180° . Para demostrarlo, podemos sumar los ángulos del «triángulo» formado por la parte del ecuador comprendida entre Kenia y Ecuador y los «segmentos de recta» que unen «puntos» de estos países con el Polo Norte. El triángulo esférico así formado tiene dos ángulos rectos en el ecuador.

Hay otras interpretaciones no estándar de los términos «punto», «recta» y «distancia» en los que son válidos los cuatro primeros axiomas de la geometría euclídea y el postulado de las paralelas es falso, aunque por otro motivo: por un punto exterior a una recta dada hay *más* de una paralela. En estos modelos (que podrían consistir, por ejemplo, en superficies en forma de silla de montar) la suma de los ángulos de un triángulo es menor que 180° . El matemático alemán Bernhard Riemann concibió superficies más generales todavía y geometrías en las que el concepto de distancia varía de un punto a otro de un modo parecido a como le ocurre a un viajero que se mueve por un terreno muy irregular y accidentado.



Si en esta superficie en forma de silla de montar interpretamos «línea recta» como la distancia más corta entre dos puntos, son válidos todos los axiomas de la geometría euclídea excepto el postulado de las paralelas. Por P pasa más de una paralela a I



Los ángulos de un triángulo sobre esta superficie suman menos de 180°

Cualquier modelo que, por la razón que sea, no cumpla el postulado de las paralelas se dice que es un modelo de geometría no euclídea. Cada una de las geometrías presentadas es un conjunto consistente de proposiciones (exactamente igual que las constituciones de distintas naciones son diferentes conjuntos de leyes consistentes entre sí). Cuál de ellas es la verdadera en el mundo real es una cuestión que depende de qué significado físico demos a los términos «punto» y «recta», y se trata de una cuestión empírica que se ha de dilucidar mediante la observación y no por proclamas de salón. Localmente al

menos, el espacio parece tan euclídeo como un campo de maíz de Iowa, pero como ya han descubierto los partidarios de la tierra plana de cualquier parte del mundo, es peligroso extrapolar demasiado lejos la propia estrechez de miras. Si tomamos la trayectoria de un rayo de luz como interpretación de línea recta, obtenemos una geometría física no euclídea.

Para terminar, me gusta pensar en el descubrimiento de la geometría no euclídea como una especie de chiste matemático —chiste que Kant no entendió—. Muchos acertijos y chistes son de la forma «¿Qué tiene esta propiedad, aquélla y la de más allá?». Al oírlos, la respuesta que se le ocurre a uno inmediatamente es completamente distinta de la interpretación inesperada de las condiciones que constituye la esencia del chiste. Así ocurre con la geometría no euclídea. En vez de «¿Qué es negro, blanco y rojo por todas partes?» tenemos «¿Qué cumple los primeros cuatro axiomas de Euclides?». La nueva esencia de este chiste nos la dieron Bolyai, Lobachevski y Gauss: grandes humoristas del Club Universo.

Gödel y su teorema

El lógico matemático Kurt Gödel fue uno de los gigantes intelectuales del siglo XX y, en el supuesto de que la especie se conserve, probablemente será una de las pocas figuras contemporáneas recordadas dentro de 1.000 años. Y aunque recientemente se hayan publicado algunos libros sobre él, esta opinión no es consecuencia de una campaña de promoción ni de una moda incipiente (aunque resulte infinitesimalmente más aceptable por el parecido entre las palabras «God», «Gödel» y «Godot»). Tampoco se trata de un caso de autocomplacencia por parte de los matemáticos, a pesar de que en todas las disciplinas sea corriente alentar una cierta miopía profesional. Sencillamente es verdad.

¿Quién fue Kurt Gödel? Su perfil biográfico no es nada complicado. Nació en 1906 en Brünn (en la actual Checoslovaquia), fue a la Universidad de Viena en 1924, donde permaneció hasta que en 1939 emigró a Estados Unidos. Vivió en Princeton, Nueva Jersey, y trabajó en el Instituto de Estudios Avanzados hasta su muerte. Durante los años treinta y principios de los cuarenta, descubrió unos resultados de lógica matemática que revolucionaron esa rama del saber. Su investigación también iluminó importantes áreas afines en la matemática, la informática y la filosofía.

El resultado más famoso, el llamado primer teorema de incompletitud, demuestra que cualquier sistema matemático formal que contenga un mínimo de aritmética es incompleto: siempre habrá enunciados que no serán demostrables ni refutables dentro del sistema, independientemente de lo elaborado que sea éste. Nadie podrá nunca escribir una lista de axiomas y pretender con razón que toda la matemática se deduce de ellos (tanto si esos axiomas ocupan todo un fajo de papel, como si ocupan una biblioteca entera

de un millón de libros, o un trillón de chips de silicio para el que hubiera hecho falta toda la arena del Sahara). Haciendo una distinción rigurosa entre los enunciados del sistema y los metaenunciados acerca del sistema, empleando ingeniosas definiciones recurrentes y asignando códigos numéricos a los enunciados de la aritmética, Gödel pudo construir un enunciado aritmético que «dice» de sí mismo que es indemostrable y con él establecer su teorema. (No me extrañaría que Boris Pasternak tuviera en mente el teorema de Gödel cuando escribió: «Lo que está asentado, clasificado y demostrado nunca será suficiente para abarcar toda la verdad»).

Una demostración alternativa del teorema, debida al informático norteamericano Gregory Chaitin, emplea conceptos de la teoría de la complejidad (véase la entrada sobre *La complejidad*). En este caso, la proposición aritmética indemostrable «dice», por medio de un código numérico, que cierta sucesión aleatoria de bits tiene una complejidad mayor que la del sistema formal dado. Por consideraciones en el metanivel del sistema, se sabe que este hecho es verdadero, pero para que la proposición sea demostrable en el sistema, éste tendría que producir una sucesión de bits que tuviera una complejidad mayor que la suya propia. Y, por la definición de complejidad, esto es imposible.

Todo ello guarda relación con la llamada paradoja de Berry. Dice así: «Encontrar el menor número entero tal que para definirlo haga falta un número mayor de palabras que las que componen esta frase». Algunos ejemplos como el número de mis cabellos, el número de configuraciones posibles del cubo de Rubik o la velocidad de la luz en milímetros por siglo definen cada una de ellas un número entero concreto con menos palabras que la frase dada. El carácter paradójico de la frase de Berry se manifiesta cuando uno se percata de que ella misma define un número entero particular que, por definición, se especifica con demasiado pocas palabras. Aunque las dos demostraciones del teorema de Gödel aprovechan conocidas paradojas, la paradoja del mentiroso, en el caso de la demostración estándar, y la de Berry en la demostración de Chaitin, la teoría de la incompletitud no es en ella misma una paradoja en absoluto. Aunque parezca extraño, se trata de matemática legítima y sin problemas.

Debería mencionarse también que, a pesar de los arduos esfuerzos por

demostrar lo contrario, el teorema no señala ninguna división fundamental entre cerebros y máquinas. Ambos están sujetos a limitaciones y condicionantes que, al menos en principio, son muy similares. A una máquina podemos incluso darle la capacidad de «mirarse desde fuera», formalizando su metalenguaje y, si hace falta, su metametalenguaje.

Los otros resultados fundamentales de Gödel tratan, entre otras cosas, del intuicionismo, la demostrabilidad y la consistencia de la matemática, y las funciones recurrentes. Más tarde, trabajó también en cosmología. Usando muchas de las ideas y elaboraciones de su primer teorema, su segundo teorema de incompletitud establece que ningún sistema matemático razonable puede demostrar su propia consistencia. Lo único que podemos hacer es suponérsela; no podemos demostrarla sin recurrir a hipótesis más fuertes que la de la propia consistencia.

Gödel tuvo una vida personal ascética, y las únicas evidencias externas de sus emociones o su personalidad son su largo matrimonio y sus periódicas depresiones, por las que tuvo que ser hospitalizado en varias ocasiones. En su juventud era algo menos solitario y se relacionó periféricamente con el Círculo de Viena, aunque no simpatizara con su positivismo. Pero, aparte de eso y su posterior amistad con Einstein en Princeton, tuvo pocos contactos sociales con sus contemporáneos. Sus relaciones con otros matemáticos se limitaban ante todo a los artículos, correspondencia y conversaciones telefónicas. Gödel nunca poseyó el compromiso apasionado de Bertrand Russell ni el fuerte sentido del humor de Einstein.

Tuvo, no obstante, otros focos de interés intelectual. Su trabajo le llevó a la convicción de que los números existían en algún dominio independiente del hombre, y que la explicación del espíritu no podía ser mecanicista, pues estaba separado de la materia y no podía reducirse a ella. Procedente de un ambiente luterano, Gödel no fue religioso en un sentido convencional, pero siempre mantuvo su teísmo y sostenía la posibilidad de una teología racional. Llegó a intentar la construcción de una variante del argumento ontológico medieval, según el cual la existencia de Dios es en cierto modo una consecuencia de nuestra capacidad de conceptualizarlo. Fue verdaderamente un gran lógico y tuvo que conocer momentos de intensa alegría intelectual. Sin embargo, un sensualista ordinario como yo no puede dejar de desear que

hubiera sido un poquito más feliz en un sentido visceral: una salud mejor, un hijo, una aventura sentimental, algo físico.

Gödel murió en Princeton el 14 de enero de 1978, según el certificado de defunción, de «malnutrición e inanición» ocasionadas por «trastornos personales».

Grupos y álgebra abstracta

El álgebra abstracta y la geometría moderna (véase la entrada sobre *Geometría no euclídea*) fueron productos del siglo XIX y ambas contribuyeron a cambiar nuestra idea de la naturaleza de la matemática. La matemática dejó de ser concebida como algo relacionado exclusivamente con verdades externas a ella y se la empezó a aceptar simplemente como un modo de deducir las consecuencias que se derivan de varios conjuntos de axiomas. Como dijo Bertrand Russell: «La matemática pura consiste únicamente en afirmaciones tales como, si de algo se puede decir tal cosa y tal otra, entonces para este algo se cumple esto y aquello». Uno de los «algos» más importantes del álgebra abstracta es un grupo.

No se trata de una cuadrilla de matemáticos con instintos gregarios. Un grupo matemático es un tipo importante de estructura algebraica abstracta. Como los grupos son abstractos, presentaré algunos ejemplos antes de dar su definición. Consideremos primero el conjunto de los números enteros, los positivos, los negativos y el cero, y la operación de la suma. Fijémonos en algunas cosas que a primera vista pueden parecer triviales: la suma de dos números cualesquiera es otro número; vale la igualdad $(3 + 9) + 11 = 3 + (9 + 11)$ y, en general, el resultado de varias sumas seguidas no depende de cómo las asociemos; existe un número, el 0, tal que, para cualquier número X , $X + 0 = X$; para cualquier número X , existe otro número que sumado a él da 0: $6 + (-6) = 0$, $(-118) + 118 = 0$, etc.

O consideremos los doce objetos $0, 1, 2, \dots, 10$ y 11 , y la operación que consiste en tomar la suma de los números si ésta es menor o igual que 11 o, en caso contrario, tomar el resto de dividir dicha suma por 12 (la operación se suele llamar suma módulo 12). Así $8 + 7 = 3$, y $(6 + 5) + 9 = 8$. Es fácil

comprobar que las tres primeras propiedades citadas más arriba también son válidas para este conjunto de objetos y esta operación. La cuarta propiedad también se cumple aunque no es tan evidente. ¿Qué hay que sumar a 7 para que dé 0? No hay un -7 , pero sí hay un 5 y $7 + 5 = 0$. Puede comprobarse que, para cada objeto, hay un «inverso» que, sumado a él, da 0.

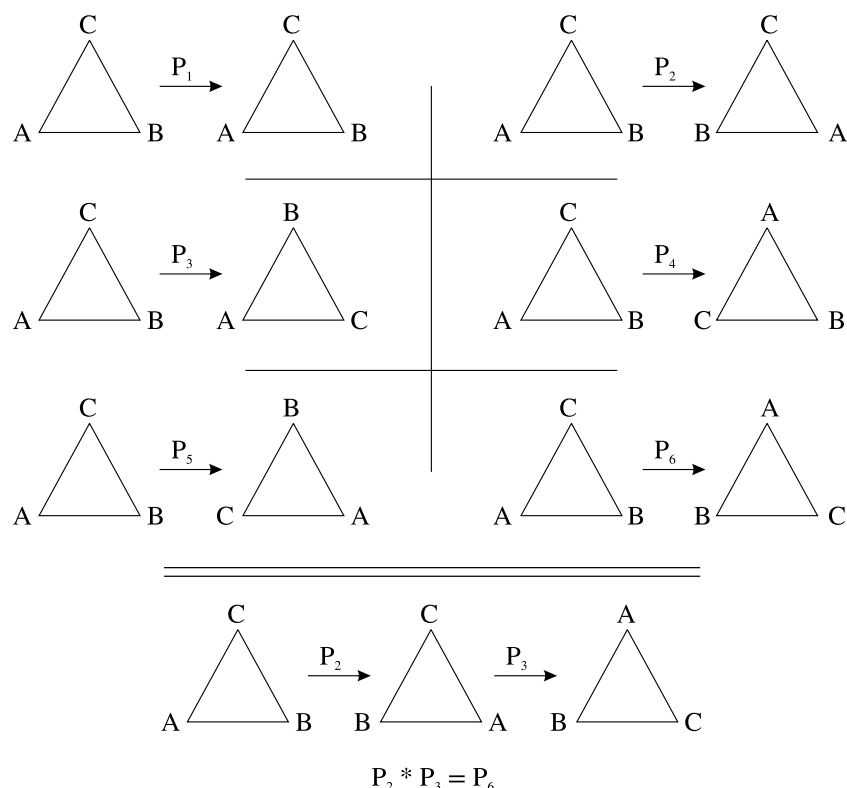
Sin dejar los números, considérense los cuatro objetos 1, 2, 3 y 4, pero esta vez la operación será como el producto, excepto que si el producto es mayor que 4, lo sustituiremos por el resto de dividir por 5 (multiplicación módulo 5). Así, $4 \times 3 = 2$. Como en los dos casos anteriores valen también las cuatro propiedades. El producto de dos objetos es otro objeto; vale la igualdad $(3 \times 2) \times 4 = 3 \times (2 \times 4)$ y en general el resultado global no depende de cómo asociemos los productos; hay un objeto, el 1 en este caso, tal que $1 \times X = X$ cualquiera que sea X ; para cada objeto X existe otro objeto que al multiplicarlo por X da 1. Para 2 es 3, pues $2 \times 3 = 1$, mientras que para 4 es él mismo, pues $4 \times 4 = 1$.

Los tres conjuntos anteriores con sus respectivas operaciones son grupos, pero retrasaré un poco más la definición general para mostrar que los grupos no tienen por qué estar formados por números o por objetos numéricos. Los elementos del próximo grupo son permutaciones (o reordenaciones) de tres objetos que, con un alarde de imaginación, llamaré A, B y C. Siguiendo con las notaciones imaginativas, denotaré la primera permutación por P_1 . No permuta nada, sino que deja A, B y C en el lugar en que estaban. P_1 se llama permutación «identidad» y juega un papel similar al que tenían el 0, el 0 y el 1 en los ejemplos anteriores, los elementos identidad de sus respectivos grupos. La siguiente permutación, P_2 , cambia A y B, pero deja C en su lugar. P_3 cambia B y C, y deja A inmóvil, mientras que P_4 cambia A y C dejando B en su lugar. P_5 pone A, B y C en el orden C, A y B, mientras que P_6 permuta A, B y C y los deja como B, C y A. Además, estas P actúan sobre cualquier ordenación de A, B y C a la que las apliquemos, de modo que, por ejemplo, el efecto de P_4 sobre B, A, C es la nueva ordenación C, A, B, resultante de intercambiar los elementos de las posiciones primera y tercera.

Para comprobar la validez de las cuatro propiedades necesitamos una operación definida en este conjunto de permutaciones $P_1, P_2, \dots$ y P_6 . La

operación será la «aplicación sucesiva»: se aplica una permutación al *resultado* de la otra y se obtiene una tercera permutación «producto» de las dos dadas. ¿Cuál sería, por ejemplo, el producto $P_2 * P_3$? Para determinarlo, sigamos los movimientos de A al realizar el producto. Como P_2 intercambia A y B, A irá a parar a la posición de B después de la primera permutación. Luego P_3 intercambia las posiciones de B y C, dejando A en el lugar de C. Sigamos a continuación la pista de B. Primero P_2 la deja en el lugar de A, y P_3 no la mueve de ahí por cuanto sólo intercambia las posiciones de B y C. En cuanto a C, P_2 no la toca y luego P_3 la pone en el lugar de B. Así pues, $P_2 * P_3$ reordena A, B y C como B, C y A, que es el mismo efecto de aplicar P_6 . Por tanto, $P_2 * P_3 = P_6$.

Quizá quiera comprobar algunos productos más de este artificio matemático. (O quizá tampoco quiera por esta vez). Por ejemplo, $P_2 * P_2 = P_1$, o $P_4 * P_6 = P_2$. De lo que se trata es de no abrumarle con el cálculo (hay artilugios de notación que nos pueden servir para ello) pero demostrarle que estas seis permutaciones con la operación de la aplicación sucesiva, $*$, satisfacen las cuatro propiedades anteriores: el producto de dos permutaciones cualesquiera es otra permutación; el resultado final de varios productos no se ve afectado por el orden en que los realicemos, como en $P_i * (P_j * P_k) = (P_i * P_j) * P_k$; hay una permutación, P_1 , tal que $P_1 * P_i = P_i$, para cualquier permutación P_i ; y para cualquier permutación P_i , existe otra permutación P_j tal que $P_i * P_j = P_1$.



Las permutaciones de A, B y C se pueden interpretar como las reflexiones y rotaciones de un triángulo y la operación $*$ como la realización sucesiva de dichos movimientos

Como ya debía esperar, la definición formal de grupo es la siguiente: cualquier conjunto con una operación definida en él y que satisfaga las cuatro propiedades anteriores. Hay un sinnúmero de ejemplos de grupos, muchos de ellos geométricos. Si hacemos corresponder cada permutación de los vértices a un movimiento de un triángulo, el grupo anterior de las permutaciones de tres elementos se puede interpretar también como un conjunto de reflexiones y rotaciones del mismo, haciendo corresponder cada permutación de los vértices a un movimiento del triángulo. (Cuando, como en este caso, tenemos dos grupos que son en esencia el mismo —elementos identidad correspondientes y operaciones correspondientes— y sólo difieren en los nombres de sus elementos y de sus operaciones, decimos que los grupos son isomorfos). Los grupos aparecen también en la teoría de los nudos y son muy útiles en la clasificación y el análisis de los nudos y las trenzas. Tienen un

papel importante en muchas otras ramas de la matemática, en cristalografía y en el estudio de los quarks y la mecánica cuántica. En la mayoría de casos, los elementos del grupo son una acción de algún tipo —una permutación, una flexión o alguna clase de función—. Hasta las distintas contorsiones de las caras de un cubo de Rubik forman un grupo.

Lo que salva a los grupos de ser una simple taxonomía matemática (definición y clasificación sin profundizar mucho más) es que se han demostrado muchos teoremas, y muy potentes, que se refieren a los grupos, subgrupos, grupos cociente y su relación con otras estructuras abstractas. Por poner un ejemplo: si G es cualquier grupo finito y H es un subgrupo cualquiera de G , entonces el número de elementos de H es un divisor del número de elementos de G . La potencia de aislar estas estructuras abstractas y demostrar teoremas relativos a ellas procede en gran parte de un hecho simple. Una vez se ha demostrado un teorema acerca de un grupo abstracto, es válido para cualquier representante del grupo, por dispar que sea. Concentrándonos en la estructura abstracta y no en los aspectos específicos, somos capaces de ver el bosque matemático en vez de los árboles concretos.

Y esto es, en fin, lo esencial del álgebra abstracta. Exactamente igual que el álgebra elemental se sirve de las variables para simbolizar números y estudiar sus propiedades, las teorías de grupos, anillos, cuerpos, espacios vectoriales y otras estructuras algebraicas llevan la abstracción mucho más allá. En estas teorías los símbolos representan conjuntos, operaciones en conjuntos, estructuras, funciones entre estructuras, etc., todo ello con el objeto de demostrar teoremas y proposiciones generales.

El resto de la cita de Russell es un final apropiado para la ocasión. «Es esencial no discutir si la proposición es realmente verdadera, ni hablar de qué es el algo acerca del cual se supone la verdad de dicha proposición [...]. Si nuestra hipótesis se refiere a algo y no a alguna o algunas cosas concretas, entonces nuestras deducciones son matemática. Así pues, la matemática se podría definir como la materia en la que nunca sabemos de qué estamos hablando, ni de si lo que estamos diciendo es cierto».

Aunque el hecho de que por todas partes haya personas que no saben de qué hablan ni si lo que están diciendo es cierto pudiera hacer pensar que el genio matemático es exuberante, la cita es un resumen sucinto, aunque

exagerado, del enfoque axiomático formal de la matemática.

Humor y matemática

Aquellas personas cuya formación escolar en la asignatura de lengua se haya reducido exclusivamente al estudio de la ortografía probablemente no sepan apreciar demasiado la literatura. Los estudiantes de bellas artes cuya preparación haya consistido principalmente en dibujar sólidos rectangulares seguramente serán impermeables a la emoción y la belleza de la pintura. Asimismo, a las personas cuya educación matemática se haya limitado a dominar las técnicas del cálculo o a cursos de rutinas enseñadas maquinalmente les parecerá extraño que hable de humor en la matemática.

Empecé a interesarme por la relación entre el humor y las matemáticas cuando me di cuenta de que, con frecuencia, los matemáticos tenían un sentido del humor característico. Quizá debido a su formación, tendían a interpretarlo todo al pie de la letra y, a menudo, sus interpretaciones literales no tenían nada que ver con las convencionales. Como la incongruencia, sea de la clase que sea, es una condición necesaria para el humor, su literalidad también era a menudo divertida. «Depositar aquí los desperdicios», por ejemplo, significa dejar las cosas en el suelo, de lo contrario pierden su condición de desperdicios. Si los ponemos en el cubo de la basura, dejan de ser desperdicios para convertirse en basura. La señal de «Bajen la voz» en la biblioteca no significa que haya que conversar bajo la mesa ni que haya que hacerlo con voz de barítono. El presentador de un desfile de bellezas que dice efusivamente, «Cada chica es más bonita que la siguiente» literalmente quiere decir que cada vez son más feas. O como: «Daría el brazo derecho por ser ambidextro», que sólo es divertido si se toma al pie de la letra.

Estos chistes un tanto pueriles son típicos de los matemáticos (o más

exactamente de algunos matemáticos, entre los cuales me cuento). Son típicas también la curiosidad por los juegos de palabras, las autoreferencias (especialmente entre mis colegas lógicos), los modelos no convencionales de situaciones, las inversiones, la reducción al absurdo, la iteración y varios otros conceptos lógicos y cuasimatemáticos. Tanto la matemática como el humor son formas de juego intelectual (véase el final de la entrada sobre *Geometría no euclídea*). La matemática está en el lado intelectual y el humor en el lado juguetón de un continuo que se extiende de la una al otro pasando por los acertijos, las paradojas y los rompecabezas. La orden que daba Lewis Carroll a los escritores de cartas, de que lamieran el sobre en vez del sello, es marginalmente divertida y marginalmente matemática (inversión de una relación lógica) y es por tanto un habitante representativo de esta zona intermedia.

Además, tanto la matemática como el humor son combinatorios, en su separar y reunir ideas por pura diversión (yuxtaponiéndolas, generalizándolas, interpolándolas o invirtiéndolas). En el humor esto es corriente, pero este mismo aspecto de las matemáticas no es tan conocido, quizá porque hace falta tener conocimientos matemáticos previos para poder jugar con ellos. Se necesita un cierto grado de sofisticación matemática para darse cuenta, por ejemplo, de que el anudado de trenzas guarda una relación sorprendentemente graciosa con alguna idea de otra rama completamente distinta de la matemática, las simetrías de una figura geométrica, pongamos por caso.

La ingenuidad y el ingenio son distintivos tanto del humor como de la matemática, al igual que una economía espartana en lo que respecta a la expresión. La prolijidad es tan antitética con la matemática pura como normalmente lo es con el buen humor. A riesgo de ser yo mismo prolijo, señalaré que la belleza de una demostración matemática depende de su elegancia y su brevedad. (Véanse las demostraciones sobre la infinitud de los números primos, de la irracionalidad de la raíz cuadrada de 2 o la innumerabilidad de los números reales). Una demostración torpe introduce consideraciones extrañas y acaba siendo tortuosa y redundante. Del mismo modo, un chiste pierde toda su gracia si se cuenta mal, se explica con demasiados detalles o se basa en analogías forzadas. Se puede perdonar que

un chiste sea de mal gusto, pero ha de ser agudo: AIXELSID.

Los patrones, las reglas, las estructuras y la lógica son componentes esenciales de la matemática y el humor, aunque no de la misma manera. En el humor, la lógica es a menudo perversa y se basa en la autoreferencia. Es el caso, por ejemplo, del votante que, cuando un encuestador le pregunta cuáles son las razones de la ignorancia y la apatía del público norteamericano, contesta, «Ni lo sé ni me importa». Los patrones y las estructuras utilizados normalmente se distorsionan y se confunden, como aquel hombre al que le cuentan que el aceite de oliva se obtiene exprimiendo aceitunas y el aceite de coco exprimiendo cocos, y entonces empieza cavilar sobre cómo debe obtenerse el aceite de niños. Y las reglas pertinentes se toman con doble sentido, como en el caso de los dos sacerdotes anglicanos que están hablando del lamentable estado de la moralidad sexual. «Yo nunca me acosté con mi mujer antes de casamos», dice uno con aire santurrón, «¿y usted?». «No estoy seguro», contesta el otro. «¿Cómo se llamaba de soltera?».

Pero estos equívocos y falsas conclusiones no son aleatorios, y han de tener sentido en algún nivel (verbigracia, las lámparas de flash solares). Para entender qué tiene de incongruente una cierta historieta, y así «coger el chiste», es esencial entender la lógica verdadera, y los patrones, reglas o estructuras correctos. Y a la inversa, hay que ser consciente de la elegancia, el vigor y la fuerza de una demostración matemática para poder apreciar plenamente qué es la matemática. Como antes también, el empleo que se da a estos entendimientos y apreciaciones en la matemática y el humor son completamente distintos. El matemático, por ejemplo, usa uno de sus gambitos preferidos, la técnica lógica de la reducción al absurdo, para demostrar proposiciones. Para demostrar S basta con suponer la negación de S y deducir de ella una contradicción o llegar a un absurdo. Los autores de comedias también usan esta técnica cuando parten de una premisa rara («Qué pasaría si...») y luego, con los consiguientes absurdos, hacen un chiste o un cuento.

Aunque los acordes emocionales no sean los mismos, los investigadores matemáticos adoptan una postura y arrogancia no muy distinta de los humoristas. La sonrisa que levanta un giro inesperado en una bella demostración es una versión enrarecida de la risa que causa el final de un

buen chiste. Terminaré con una pregunta fácil: ¿Qué pregunta contiene la palabra cantalupo sin motivo aparente?

[Nota: AIXELSID es DISLEXIA escrito al revés].

Imposibilidades: tres viejas y tres nuevas

A menudo los estudiantes se preocupan demasiado por hallar soluciones a los problemas y rara vez se dan cuenta de que algunos de los problemas más interesantes de la matemática (y de otros campos) carecen de solución. En esta categoría se enmarcan tres antiguos problemas de construcción clásicos y también tres resultados revolucionarios del siglo XX.

En las construcciones sólo se puede usar una regla sin graduar y un compás para: (1) la duplicación de un cubo, (2) la trisección de un ángulo y (3) la cuadratura del círculo. Como estas expresiones no son suficientemente explicativas (duplicar un cubo suena un poco a juego de manos), las detallaré un poco más. Dado un cierto segmento de recta, se duplica un cubo obteniendo otro segmento de recta tal que el cubo que lo tenga por arista tenga un volumen doble que el cubo cuya arista sea el segmento original. Esto no se consigue doblando sin más la longitud del segmento original, pues entonces el volumen del cubo resultante sería ocho veces mayor que el del primero, y no dos. El significado de la trisección de un ángulo es bastante claro: se trata de dividir el ángulo en tres partes iguales. El método ha de ser válido para cualquier ángulo, no sólo para algunos casos particulares. Y cuadrar un círculo significa hallar un segmento tal que el área del cuadrado que lo tenga por lado sea igual al área del círculo dado. Recordemos también que se pide un método que en principio sea exacto, no sólo aproximado.

Hasta el siglo XIX no se demostró concluyentemente la imposibilidad de estas construcciones y, hasta entonces, los matemáticos siguieron intentando hallar las soluciones. En las ecuaciones asociadas a estos problemas intervienen ya sea raíces cúbicas, o el número pi —que no satisface ninguna

ecuación algebraica (véase la entrada sobre Pi)—, y se demostró que los números construibles con regla y compás se reducían a los que se definían mediante raíces cuadradas (y raíces cuadradas de raíces cuadradas). No obstante, las matemáticas que nacieron a partir de estos intentos fallidos son muchísimo más valiosas, tanto teórica como prácticamente, que las propias construcciones. El trabajo sobre curvas analíticas, las ecuaciones cúbicas y cuárticas, la teoría de Galois y los números trascendentes, todo ello procede, al menos en parte, de estas imposibilidades.

El fenómeno del grano de arena dentro de la ostra es muy general en la matemática. El reconocimiento de la irracionalidad de la raíz cuadrada de 2 echó por los suelos la creencia pitagórica de que todas las cosas eran explicables en términos de los números enteros y sus cocientes, pero estimuló a Eudoxo, Arquímedes y otros a desarrollar una teoría de los números irracionales. Los problemas que presentaba deducir el postulado de Euclides de las paralelas a partir de los otros cuatro condujeron finalmente a un gran descubrimiento del siglo XIX (en realidad el descubrimiento fue de János Bolyai, Nikolái Lobachevski y Karl Friedrich Gauss): el de la geometría no euclídea. Las rarezas de la teoría de conjuntos de Georg Cantor en la proximidad del cambio de siglo sacudieron la comunidad matemática y fueron en parte una causa del trabajo fundamental de Bertrand Russell, Alfred North Whitehead y otros.

Sorprendentemente, tres de los descubrimientos más significativos del siglo XX toman la forma de afirmaciones de imposibilidades. El teorema de incompletitud de Kurt Gödel afirma que en cualquier sistema matemático axiomático siempre habrá proposiciones que no son ni demostrables ni refutables; en otras palabras, es imposible demostrar dentro del sistema todas las verdades acerca del sistema. Esto acabó con las esperanzas de aquellos que pensaban que todas las verdades matemáticas se podían derivar de un simple sistema axiomático. (Véase la entrada sobre Gödel).

En física, el principio de incertidumbre de Heisenberg afirma que es imposible determinar exactamente la posición y el momento de una partícula en un instante dado cualquiera, y que el producto de las incertidumbres de estas dos cantidades siempre es mayor o igual que cierta constante. Además de su impacto revolucionario en la física, el principio de incertidumbre hizo

que resultara mucho más problemática la profesión de una filosofía de la ciencia estrictamente determinista. Desgraciadamente, también ha servido de inspiración a muchos «parapsicólogos». Las comillas son para indicar mi firme convicción de que quienes practican una disciplina deberían tener primero una disciplina que practicar. Impresionados por algunas semejanzas formales vagas entre la mecánica cuántica y el comportamiento humano —la incertidumbre en la predicción, una enorme sensibilidad a la observación, el fracaso de ciertas leyes lógicas, la inconmensurabilidad o complementariedad de «puntos de vista»— estos parapsicólogos han dado un lustre científico a la creencia en la telepatía, la psicoquinesis y la precognición, sin proporcionarles ninguna base científica.

Finalmente, el teorema del economista Kenneth Arrow sobre las «funciones de la elección social» afirma que no hay ninguna manera infalible de obtener las preferencias de un grupo a partir de las preferencias individuales que garantice el cumplimiento de ciertas condiciones mínimas razonables. En otras palabras, es imposible diseñar un sistema de votación que no pueda presentar defectos graves en ninguna ocasión. Al igual que con Gödel y Heisenberg, hubo que abandonar también un ideal intelectual, en este caso la esperanza de encontrar un método universal para las elecciones sociales. (Véase la entrada sobre *Sistemas de votación*).

El reconocimiento de la imposibilidad teórica es a menudo un indicio de sofisticación intelectual. La gente y las sociedades primitivas normalmente pueden resolver todos sus problemas.

La inducción matemática

Imaginemos una escalera de mano con infinitos peldaños que llegue hasta el cielo, la azotea de Platón (sobre su famosa cueva), o donde sea. Podremos llegar a esta posición de privilegio si somos capaces de subimos al primer peldaño (o a algún otro) de la escalera y si, siempre que estamos en un peldaño dado (llamémosle K), somos capaces de subir al siguiente ($K + 1$). La formalización de esta idea es el axioma de la inducción matemática, una de las armas más potentes del arsenal matemático. (Véase también la entrada sobre *La recurrencia*). Otra metáfora ilustrativa es una hilera infinita de fichas de dominó. Si cae una, seguirá el resto, pero si no empujamos la primera o falta una ficha, entonces no caerán todas.

Consideremos a modo de ilustración la demostración de que $1 + 2 + 3 + 4 + \dots + N = [N(N + 1)]/2$. Tomamos primero un número cualquiera, pongamos el 7, y comprobamos si la fórmula es cierta para él. ¿Es verdad que $1 + 2 + 3 + 4 + 5 + 6 + 7 = (7 \times 8)/2$? La respuesta es sí, pues ambos miembros valen 28. Se podría comprobar también que la fórmula es válida para 10. Aunque estos casos y otros que pudiéramos considerar sugieren la posibilidad de que la fórmula quizá valga para todos los enteros, no la demuestran. Para hacerlo usaremos la inducción matemática. ¿Vale la fórmula para $N = 1$?, o, por emplear la analogía de la escalera, ¿podemos subimos al primer peldaño? Sustituyendo 1 en la fórmula, obtenemos $1 = (1 \times 2)/2$, que naturalmente se cumple. La fórmula vale también para $N = 2$, pues $1 + 2 = (2 \times 3)/2$.

Ya estamos pues en la escalera. Pero ¿podemos subir siempre al siguiente peldaño? Supongamos que estamos balanceándonos precariamente en el K -ésimo peldaño (para un K fijo escogido arbitrariamente) y que la fórmula es

válida para este K : $1 + 2 + 3 + 4 + \dots + K = [K(K + 1)]/2$. Para demostrar que podemos trepar al siguiente peldaño, hemos de probar su validez también para $(K + 1)$, esto es, que $1 + 2 + 3 + 4 + \dots + K + (K + 1) = [(K + 1)(K + 2)]/2$. Podemos sustituir los K primeros términos del primer miembro, $1 + 2 + 3 + 4 + \dots + K$, por su suma, que por la primera ecuación supondremos igual a $[K(K + 1)]/2$. Hecha la sustitución nos queda $[K(K + 1)]/2 + (K + 1) = [(K + 1)(K + 2)]/2$.

Nos queda por demostrar esta ecuación y, para ello, basta simplemente con desarrollar algebraicamente ambos miembros y comprobar que son efectivamente iguales. (Aquellos a los que la palabra «simplemente» les suene a broma deberían intentar seguir la lógica de los pasos precedentes, que en cualquier caso es mucho más sofisticada e interesante que el álgebra). Satisfechas ambas condiciones del axioma de inducción, alcanzamos el cielo matemático y decimos que la fórmula es válida para todos los valores enteros de N .

Comprobar la veracidad de una proposición dada para varios casos particulares quizá contribuye a hacerla verosímil, pero no constituye ni mucho menos una demostración inductiva de la misma. Podemos comprobar, por ejemplo, que para muchos valores de N la suma $(N^2 + N + 41)$ es un número primo. (Los números primos no pueden descomponerse en factores como podemos hacer por ejemplo con 35; $35 = 5 \times 7$). Para $N = 1$, la expresión $(N^2 + N + 41)$ es igual a 43; para $N = 2$, da 47; para $N = 3$, 53; para $N = 4$, 61; para $N = 5$, 71; para $N = 6$, 83; para $N = 7$, 97; para $N = 8$, 113; para $N = 9$, 131. De hecho, todos los valores de N hasta 39 dan un primo para $(N^2 + N + 41)$. Pero para $N = 40$ la proposición es falsa.

La inducción matemática puede usarse para demostrar cualquier proposición relativa a un entero arbitrario. Podemos usarla, por ejemplo, para demostrar que la suma de los ángulos de un polígono convexo (triángulo, cuadrilátero, pentágono, hexágono, etc.) es siempre igual a $(N - 2) \times 180^\circ$. Cómo llegamos a formular una de esas proposiciones es un tema muy distinto de cómo las demostramos y, de hecho, no se presta a un análisis formal. En este sentido, quiero insistir en que no hay que confundir la inducción matemática con la inducción científica, que consiste, al menos en una primera

aproximación, en inferir leyes empíricas generales a partir de ejemplos concretos. Los teoremas demostrados por inducción matemática son deductivos y ciertos, mientras que las conclusiones a que llegamos por inducción científica son, en el mejor de los casos, sólo probables.

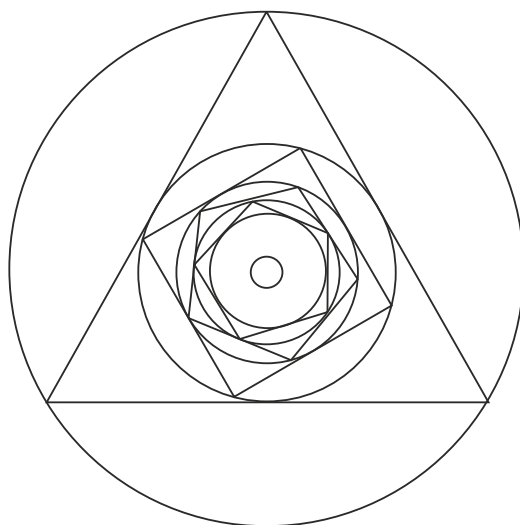
Para reforzar su confianza, puede examinar la afirmación de que $(4^N - 1)$ es divisible por 3 para cualquier valor de N . Se cumple para $N = 2$, por ejemplo, pues $(4^2 - 1)$ (que vale 15) es divisible por 3. ¿Podría demostrar por inducción el resultado general?

Y para reforzarla más todavía, demuestre que la inducción lleva verdaderamente al cielo. Demuestre que sólo le hacen falta dos requisitos para alcanzar la inmortalidad: nacer y tener garantías de que, cualquiera que sea el día, vivirá hasta el día siguiente.

[Para los interesados, la demostración de la penúltima proposición. $(4^N - 1)$ es en efecto divisible por 3. Supongamos ahora que $(4^K - 1)$ es divisible por 3 y demostremos que $(4^{(K+1)} - 1)$ también lo es. En primer lugar nos damos cuenta de que $[4^{(K+1)} - 1]$ es igual a $[(4 \times 4^K) - 1]$. Luego hacemos un poco de manipulación algebraica y, sumando y restando 4, escribimos la expresión anterior como la suma: $[4 \times (4^K - 1)] + (4 - 1)$. Como, por hipótesis $(4^K - 1)$ es divisible por 3, también lo es $[4 \times (4^K - 1)]$. Finalmente, como $(4 - 1)$ es divisible por 3, y como la suma de dos expresiones divisibles cada una por 3 lo es también, llegamos a la conclusión de que $[4^{(K+1)} - 1]$ es divisible por 3].

Límites

Tome un círculo de un metro de diámetro e inscriba en él un triángulo equilátero. A continuación, inscriba un círculo en este triángulo y luego un cuadrado dentro de este círculo más pequeño. Inscriba en el cuadrado un nuevo círculo, en el que inscribirá un pentágono regular. Continúe con estas inscripciones encajadas, alternando círculos y polígonos regulares con un lado más a cada iteración. Está claro que el área de las figuras inscritas disminuye con cada repetición, pero ¿cuál es el área final a la que se llega con esta sucesión de figuras? A primera vista parece como si tuviera que ser cero, y que el proceso concluyera en un único punto aislado. Recuerde, no obstante, que a medida que aumenta el número de lados de los polígonos, éstos se hacen cada vez más circulares y, al cabo de un cierto tiempo, el proceso se convierte, o por lo menos casi se convierte, en circunscribir un círculo dentro de otro, con una pérdida de área mínima entre cada paso y el siguiente. En cualquier caso, ya no le queda más tiempo para pensar la respuesta. El límite de este proceso es un círculo, concéntrico con el primero y con un diámetro aproximadamente 12 veces menor que el de éste.



Inscripciones encajadas que alternan círculos y polígonos regulares, cuyo número de lados aumenta en uno a cada iteración. El límite es un círculo de diámetro aproximadamente 12 veces menor que el original

Muchos otros problemas geométricos acaban en límites y, tradicionalmente, se dice que esta idea es el concepto fundamental del cálculo, la rama de la matemática que trata del cambio. Con la idea de límite podemos hacer que tenga sentido una tasa de variación instantánea, por ejemplo de la posición de un satélite en el espacio, o la suma exacta de una cantidad que varía continuamente, como la fuerza total sobre una presa inclinada. (Véase la entrada sobre *Cálculo*). Sin este concepto nos veríamos obligados a trabajar sólo con aproximaciones y promedios. Los límites representan formalmente nuestra intuición de algo que se aproxima o tiende a un valor final y aclaran la relación entre las figuras ideales y el infinito.

No obstante, discrepo del lugar principal que suele asignarse a los límites en el primer curso de cálculo, donde la principal consecuencia de su estudio es una reducción drástica de la comunidad de futuros estudiantes de matemáticas, ciencia e ingeniería. La ignorancia de una definición precisa de límite y de los otros conceptos que éste lleva asociados no fue ningún impedimento grave para los inventores del cálculo, Isaac Newton y Gottfried Wilhelm von Leibniz. Su tarea se vio coronada con el éxito, a pesar de que sólo tenían una idea intuitiva de los límites. De hecho, con el cálculo y sus otras ideas fundamentales sobre el movimiento y la gravitación, Newton dio

paso a los avances más revolucionarios de todos los tiempos en nuestra concepción del mundo físico. La ignorancia de dichas definiciones tampoco entorpeció la prolífica obra de Leonhard Euler y de sus compatriotas, la familia de matemáticos suizos Bernouilli, en el siglo XVIII. Y por la misma razón, tampoco es un obstáculo para muchos físicos e ingenieros de hoy el no ser plenamente conscientes de tales conceptos.

El problema anterior de las áreas encajadas tiene la característica bastante corriente de que para resolverlo basta con algo de ingenio y una idea somera de lo que son los límites. La opinión tradicional del papel central de los límites es verdadera, sin embargo, en el sentido siguiente. Para el progreso teórico posterior en especialidades matemáticas como las ecuaciones diferenciales, las series infinitas, el cálculo de variaciones, y el análisis real, complejo y funcional, es absolutamente necesario disponer de unos fundamentos más rigurosos que los de Newton con su «método de fluxiones» que, tomado al pie de la letra, era un disparate.

Estos fundamentos más rigurosos fueron asentados en el siglo XIX por los matemáticos Augustin Louis Cauchy, Richard Dedekind y Karl Weierstrass. Entre otras cosas, definen el límite de una sucesión numérica como el número L tal que, por pequeño que sea un número ε dado, la diferencia entre los términos de la sucesión y el número L se hace finalmente (de manera permanente) más pequeña que ε . Así, en particular, la sucesión $1/2, 3/4, 7/8, 15/16, 31/32, \dots$ tiende a 1 porque para *cualquier* número pequeño ε se puede demostrar que la diferencia entre los términos de la sucesión y el número 1 a la larga se hace (permanentemente) menor que ε .

Una vez tenemos esta definición (que es equivalente a varias otras) podemos definir el límite de una función $Y = f(X)$ para X tendiendo a un número A . Formalmente, este límite es L si, para cualquier sucesión de valores de X que tiene A por límite, la correspondiente sucesión de valores de Y (obtenidos mediante la función) tiende al límite L . Intuitivamente, Y se acerca a L tanto como queramos siempre que X esté suficientemente cerca de A . Con esta definición podemos caracterizar con precisión el concepto fundamental de derivada de una función. Definiéndola como un límite (de una cierta función cociente asociada), evitamos muchas de las críticas que levantó Newton. Éstas se centraban en la descripción de Sir Isaac de

velocidad instantánea de un objeto, como el límite informal de los cocientes de las distancias recorridas entre los tiempos transcurridos, y la subsiguiente necesidad de explicar cómo estas cantidades que se anulaban podían a la vez ser cero y no serlo. (Si aquí se ha perdido, no se preocupe; está en buena compañía).

Aunque sea un error insistir *prematuramente* en muchos puntos concretos delicados relativos a los límites, también lo es fiarlo todo a la intuición. La definición precisa de límite hace falta para aclarar qué significa el área de una región curva, el límite asintótico de una función o de una sucesión de funciones, la suma de una serie ($1 + 1/3 + 1/9 + 1/27 + \dots$) (véase la entrada sobre *Series*), y multitud de otras construcciones matemáticas más esotéricas. En un tono gracioso poco normal en él, Newton dijo en cierta ocasión que si él había visto más lejos que otros era porque se había aupado sobre los hombros de gigantes. Aupándonos en sus hombros, en los de Cauchy y en los de una miríada de otros hombros, podemos ver aún más lejos. El límite de esta sucesión de hombros sobre los que auparse es indeterminado.

Matrices y vectores

El significado latino original de la palabra «matriz» es útero. Por extensión ha pasado a significar también aquello en cuyo interior o a partir de lo cual se origina y desarrolla algo. En comparación con esto, el significado que se da en matemáticas a esa palabra es bastante estéril. Una matriz es una tabla de números ordenados en filas y columnas, y su dimensión viene dada por dos enteros que indican el número de las mismas.

Las matrices $\begin{pmatrix} 2 & 2 & 3 & -1 & 9 \\ 8 & 3 & 7 & 11 & -2 \end{pmatrix}$ y $\begin{pmatrix} 7 & 6 \\ 1 & 0 \\ 3 & 6 \end{pmatrix}$ son matrices de 2 por 5 y

de 3 por 2, respectivamente, mientras que una matriz de 1 fila y N columnas se llama generalmente vector N-dimensional. Estas definiciones no son muy apasionantes. Probablemente las disposiciones tabulares de números se conozcan desde los tiempos de los primeros contables fenicios. Lo relativamente moderno son las interpretaciones que se han dado a este sencillo instrumento de notación y a las propiedades del sistema algebraico resultante de definir operaciones entre matrices.

La aplicación matemática más corriente de las matrices está relacionada con la resolución de sistemas de ecuaciones lineales que aparecen en muchos contextos físicos y económicos. (Véase la entrada sobre *Programación lineal*). El método es parecido al que se usa en álgebra elemental. Para resolver simultáneamente $12W + 3X - 4Y + 6Z = 13$, $2W - 3Y + 2Z = 5$, $34W - 19Z = 15$, y $5W + 2X + Y - 3Z = 11$ hay que multiplicar varias de estas ecuaciones por números escogidos de modo que, al sumar o restar las ecuaciones dos a dos, las variables vayan desapareciendo sucesivamente. Después de resolver unos cuantos sistemas de esta clase, la automatización de

un procedimiento tan tedioso se plantea como una idea muy interesante, y la práctica acaba por consistir en hacer los cálculos sólo en la matriz formada por los números que aparecen en las ecuaciones (o coeficientes) y prescindiendo de ellas. La matriz de coeficientes del sistema anterior es

$\begin{pmatrix} 12 & 3 & -4 & 6 & 13 \\ 2 & 0 & -3 & 2 & 5 \\ 34 & 0 & 0 & -19 & 15 \\ 5 & 2 & 1 & -3 & 11 \end{pmatrix}$, y las distintas operaciones aritméticas que hay que

realizar con las filas de esta matriz (que corresponden a las ecuaciones) la reducen a una matriz más simple en la que figura la solución de las ecuaciones.

Además de estas operaciones con las filas y columnas de una matriz, hay operaciones con la matriz como un todo que resultan esenciales para otras aplicaciones. Para no excederme en los cálculos, consideraré primero sólo las matrices A y B : $\begin{pmatrix} 3 & -6 \\ 5 & 2 \end{pmatrix}$ y $\begin{pmatrix} 1 & 0 \\ 3 & -8 \end{pmatrix}$, respectivamente. Las matrices $(A + B)$, $(A - B)$, $5A$, $(5A - 2B)$ y $A*B$ son, respectivamente:

$$\begin{pmatrix} 4 & -6 \\ 8 & -6 \end{pmatrix}; \begin{pmatrix} 2 & -6 \\ 2 & 10 \end{pmatrix}; \begin{pmatrix} 15 & -30 \\ 25 & 10 \end{pmatrix}; \begin{pmatrix} 13 & -30 \\ 19 & 26 \end{pmatrix}; \begin{pmatrix} -15 & 48 \\ 11 & -16 \end{pmatrix}.$$

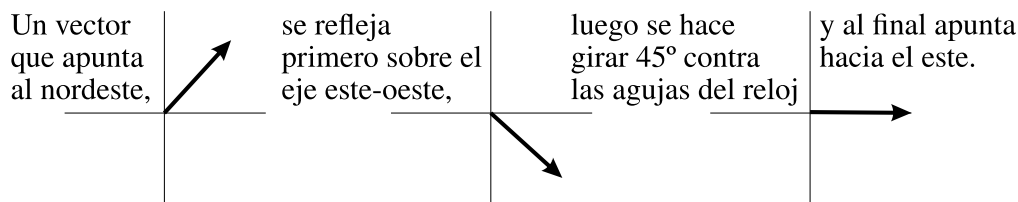
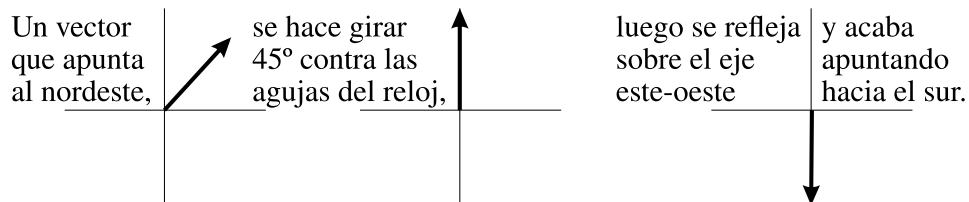
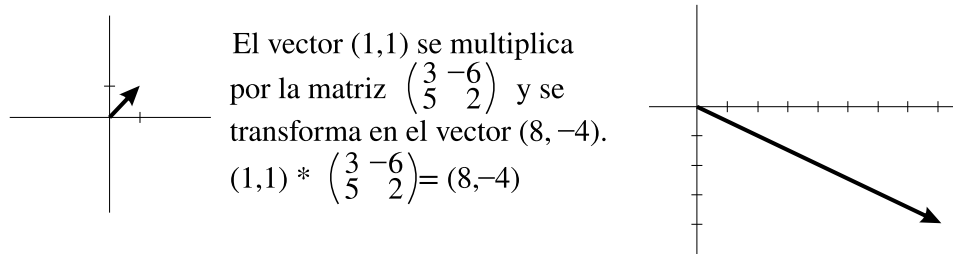
Cada elemento (o componente) de la suma $(A + B)$ se obtiene sumando los correspondientes elementos de A y de B . Para obtener $(A - B)$, se restan los elementos correspondientes (si tiene el álgebra un poco oxidada, recuerde que restar -6 equivale a sumar $+6$). Para multiplicar un número por una matriz basta con multiplicar cada elemento de ésta por dicho número. Esto explica cómo obtener $5A$ y $-2B$. Sumando estas dos se obtiene $(5A - 2B)$.

¿Y el producto $A*B$? El elemento de la primera fila y la primera columna de $A*B$ es $[(3 \times 1) + (-6 \times 3)]$, esto es, -15 ; se calcula multiplicando componente a componente los elementos de la primera fila de A por los de la primera columna de B y sumando los resultados. El elemento de la primera fila y la segunda columna de $A*B$ es 48 y se calcula multiplicando también componente a componente la primera fila de A por la segunda columna de B y sumando los resultados. Y, en general, el elemento de la i -ésima fila y la j -ésima columna de $A*B$ se obtiene multiplicando componente a componente

los elementos de la i -ésima fila de A por los de la j -ésima columna de B y sumando los productos. Aplicando este método obtenemos $A*B = \begin{pmatrix} -15 & 48 \\ 11 & -16 \end{pmatrix}$. Finalmente, la matriz $I = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix}$ se llama matriz identidad (¿útero de gemelos?), pues $C*I = I*C = C$, para cualquier matriz C.

¿Y qué? Una consecuencia puramente matemática de estas definiciones es que el conjunto de matrices forma una estructura algebraica que se llama anillo no conmutativo. Renunciaré a dar la definición de anillo (a grandes rasgos, es un conjunto con un par de operaciones definidas entre sus elementos y que cumplen ciertas propiedades) pero señalaré que «no conmutativo» significa que, contrariamente a lo que ocurre con los números, $A*B$ no tiene por qué ser igual a $B*A$. En nuestro ejemplo $A*B$ es $\begin{pmatrix} -15 & 48 \\ 11 & -16 \end{pmatrix}$ mientras que $B*A$ es $\begin{pmatrix} 3 & -6 \\ -31 & -34 \end{pmatrix}$.

Un poco de comprensión del funcionamiento de los vectores puede aclarar el por qué de esta no conmutatividad del producto de matrices. Así pues, siguiendo con el inexorablemente didáctico cursillo de esta entrada, digo que los vectores N-dimensionales (las matrices 1 por N) se emplean para representar magnitudes que se caracterizan, aparte de por su intensidad (como las temperaturas o los pesos), por su dirección y sentido (como las fuerzas y los campos electromagnéticos). De hecho, suele ser útil pensar en los vectores como si fueran flechas de una longitud conveniente apuntando en la dirección apropiada.



Estas transformaciones no conmutan, y tampoco lo harán las matrices que las representan

La velocidad es una magnitud vectorial típica. Una velocidad del viento de 10 kilómetros por hora se podría indicar por cualquiera de los vectores $(0, 10)$, $(10, 0)$, $(0, -10)$ ó $(-10, 0)$, dependiendo de que la dirección del viento fuera, respectivamente, *hacia* el norte, el este, el sur o el oeste. En cada caso, el primer número es la componente de la velocidad en la dirección este-oeste y el segundo la componente en la dirección norte-sur. El vector $(-7, 1)$ tiene una longitud 10 [determinada por el teorema de Pitágoras: $(-7, 1)^2 + 7, 1^2 = 10^2$] y por tanto se puede considerar también que indica una velocidad de 10 kilómetros por hora, pero en una dirección comprendida entre el oeste y el norte (conocida como noroeste). El vector $(9, 4)$ indica también una velocidad de 10 (pues $9, 4^2 + 3, 4^2 = 10^2$), pero esta vez en una dirección 20° al norte del este. Un viento cuyas componentes de la velocidad fueran $(28, 2)$, tres veces mayores que las de $(9, 4)$, soplaría en la misma dirección pero a 30 kilómetros por hora.

En general, se usan los vectores para representar magnitudes cuya especificación precisa de dos o más dimensiones y no hace falta que representen nada físico. Tratando de restaurantes podría ser útil introducir vectores pentadimensionales para indicar una clasificación numérica de los mismos que atendiera a cinco criterios distintos.

Comoquiera que interpretemos los vectores, las matrices se pueden considerar como transformaciones de los mismos; al multiplicarlo por una matriz (definido precisamente como el producto de dos matrices) un vector se transforma en otro vector. El vector $(1, 1)$ se transforma en $(8, -4)$ al multiplicarlo por $\begin{pmatrix} 3 & -6 \\ 5 & 2 \end{pmatrix}$ porque $(1, 1) * \begin{pmatrix} 3 & -6 \\ 5 & 2 \end{pmatrix} = (8, -4)$. Estas «transformaciones lineales» alargan, giran y reflejan los vectores (aunque si los vectores no denotan cantidades físicas, estos alargamientos, giros y reflexiones no son más que un modo de hablar) y siempre se pueden representar por matrices.

Si se realizan una tras otra, estas transformaciones de vectores no tienen por qué conmutar. Por ejemplo, una rotación de un vector seguida de una reflexión del vector resultante no da siempre lo mismo que la reflexión seguida de la rotación. (Para verlo, piénsese en un vector que apunta hacia el nordeste y es girado 45° en sentido contrario a las agujas del reloj, de manera que quede apuntando hacia el norte. Si le hacemos sufrir una reflexión sobre el eje este-oeste, al final apuntará hacia el sur. Si en vez de ello hacemos primero la reflexión sobre el eje este-oeste y luego la rotación de 45° contra las agujas del reloj, el vector acaba apuntando hacia el este). Así pues, las matrices que representan estas rotaciones, reflexiones y otras transformaciones tampoco tienen por qué conmutar: $A*B$ no siempre es igual a $B*A$.

Las matrices y los vectores juegan un papel principal en el álgebra lineal y también en muchas otras áreas de la matemática aplicada. La falta de conmutatividad explica en parte el porqué de la importancia de las matrices en mecánica cuántica, donde el orden en que se realizan dos medidas afecta el resultado final. Los tensores, que son una generalización natural de las matrices, constituyen uno de los principales ingredientes de la formulación matemática de la teoría de la relatividad general. Quizá, después de todo, la

etimología de la palabra «matriz» no sea tan inapropiada.

Media, mediana y moda

El estudiante de cuarto grado observa que la mitad de los adultos del mundo son hombres y la otra mitad mujeres y saca de ello la conclusión de que el adulto medio tiene un pecho y un testículo. Un agente inmobiliario le informa de que el precio medio de una casa en cierto barrio es de 40.000.000 ptas. y de ello deduce que en dicha vecindad hay muchas casas sobre este precio. Un vendedor dice que la mediana de las comisiones de sus nueve ventas de hoy es de 8.000 ptas. y sugiere que ganó 72.000 ptas. con esas ventas. El dueño del restaurante dice que la moda, o lo más corriente, de las cuchipandas en sus fiestas es 120.000 ptas. e insinúa que la mitad de sus clientes gastan más. Un agente de bolsa afirma que su inversión valdrá millones pero, por los cálculos que usted hace, piensa que cientos es más ajustado.

De la única de las cinco afirmaciones de la que podemos estar seguros es la del estudiante de cuarto grado. La media, la mediana y la moda son «indicadores medios» o medidas de la tendencia central, unos números que pretenden dar una idea de lo que es típico y corriente en una situación dada, pero no siempre es así. Como sus valores relativos pueden variar considerablemente, es importante conocer sus definiciones. (Véase también la entrada sobre *Estadística: dos teoremas*).

La media de un conjunto de números es lo que normalmente se conoce como promedio (o media aritmética) de dichos números. Para encontrar la media de N números basta con hallar su suma y dividirla por N . La definición es fácil y conocida, pero muchas de las inferencias que la gente saca de ella carecen de fundamento. Por ejemplo, el barrio referido anteriormente puede tener muy pocas casas de más de 40.000.000 ptas.; quizás hay en él unas

pocas grandes mansiones carísimas rodeadas de un enjambre de casas modestas.

Por contra, la mediana de un conjunto de números es el número situado en el centro del conjunto. Para hallarla basta con alinear los números en orden creciente; la mediana es el número del medio (o la semisuma de los dos números centrales, si el conjunto de números es par). Así, la mediana del conjunto 8, 23, 9, 23, 3, 57, 19, 34, 12, 11, 18, 95 y 48 se determina ordenándolos como 3, 8, 9, 11, 12, 18, 19, 23, 23, 34, 48, 57 y 95 y observando que 19 es el que ocupa el lugar central de la lista, por lo cual es la mediana de este pequeño conjunto de números, cuya media es, por cierto, 27,7. (En grandes colecciones de números la mediana se llama a veces 50-ésimo percentil, indicando con ello que es mayor que el 50% de los números de la colección. Análogamente, cuando se dice que un número está en el 93-ésimo percentil se está diciendo que es mayor que el 93% de los números).

El vendedor del ejemplo anterior podría haber ganado millones de pesetas en comisiones por sus nueve ventas de hoy, con una mediana de 8.000 ptas.; quizá seis de las ventas le supusieron una comisión de 8.000 ptas. cada una, mientras que ganó 500.000 ptas. por cada una de las tres restantes. La mediana de las comisiones sería, tal como dijo, 8.000 ptas. pero el total de sus comisiones sería mucho más que las 72.000 ptas. que él sugería haber ganado. La comisión media en este caso sería de 172.000 ptas.

Otro número, que a veces conduce a más equívocos, es la *moda* de un conjunto de datos. No es más que el dato que aparece con más frecuencia y no tiene por qué estar cerca de la media ni de la mediana del conjunto. El dueño del restaurante que decía que 120.000 ptas. era la moda de las cuchipandas quizás había tenido los siguientes pedidos en aquel mes: 40.000 ptas., 80.000 ptas., 80.000 ptas., 120.000 ptas., 80.000 ptas., 120.000 ptas., 20.000 ptas., 120.000 ptas., 20.000 ptas., 40.000 ptas., 20.000 ptas., 120.000 ptas., 40.000 ptas. La moda es efectivamente 120.000 ptas., pero la mediana es 80.000 ptas. y la media sólo 69.200 ptas.

Un ejemplo algo más sofisticado y truculento de la diferencia entre la media y la moda de una cantidad lo tenemos en el hombre que invierte 100.000 ptas. en un valor volátil que cada año sube un 60% o baja un 40%

con la misma probabilidad. Estipula que el valor ha de pasar en herencia a su nieta, quien no ha de venderlo hasta dentro de 100 años, y se pregunta cuánto le darán por él. Esta cantidad depende del número de años en que el valor haya experimentado un alza, y la media matemática, que en contextos probabilísticos se llama también su esperanza matemática, es la friolera de 1.378.000.000 ptas. Sin embargo, la moda o valor más probable de su herencia es sólo la miseria de 13.000 ptas.

La explicación de esta gran diferencia es que las rentas astronómicas correspondientes a muchos años de alza del 60% sesgan la media hacia arriba, mientras que las pérdidas correspondientes a muchos años de un 40% de baja están acotadas inferiormente por 0 ptas. El problema es una versión contemporánea de la llamada paradoja de San Petersburgo. [Para los interesados en más detalles: el valor aumenta una media del 10% anual (el promedio entre +60% y -40%). Así, al cabo de 100 años, la media o esperanza matemática de la inversión es $100.000 \text{ ptas.} \times (1,10)^{100}$, que es 1.378.000.000 ptas. Pero por otra parte, el resultado más probable es que el valor experimente un alza en exactamente 50 de los 100 años. Por tanto, la moda es $100.000 \text{ ptas.} \times (1,6)^{50} \times (0,6)^{50}$, que es 13.000 ptas. La esperanza matemática no siempre coincide con el valor que se espera].

Corrientemente, la esperanza matemática de una cantidad se calcula multiplicando sus posibles valores por las probabilidades correspondientes a los mismos y sumando estos productos. Consideremos a modo de ilustración una compañía de seguros domésticos que tiene sus buenas razones para creer que en promedio, cada año, por cada 10.000 de sus pólizas recibirá una reclamación de 40.000.000 ptas.; una de cada 1.000 reclamará una indemnización de 5.000.000 ptas.; una de cada 50 reclamará 200.000 ptas., y el resto no dará ningún problema. La compañía de seguros quiere saber cuál es su desembolso medio por póliza (para saber qué primas ha de cobrar) y la respuesta es la esperanza matemática. En este caso: $(40.000.000 \text{ ptas.} \times 1/10.000) + (5.000.000 \text{ ptas.} \times 1/1.000) + (200.000 \text{ ptas.} \times 1/50) + (0 \text{ ptas.} \times 9.789/10.000) = 4.000 + 5.000 + 4.000 + 0 = 13.000 \text{ ptas.}$

La comprensión de estas distintas medidas de la tendencia central hará un 36,17% menos probable que la persona media sea víctima de los usos

engañosos de estas cantidades por parte de los agentes de la propiedad inmobiliaria, vendedores, agentes de bolsa y dueños de restaurantes. Naturalmente, esta misma persona ya habrá sido víctima de un caso grave de hermafroditismo y, por tanto, él o ella tendrá probablemente otras preocupaciones.

El método de simulación de Montecarlo

Se sabe que un determinado jugador de baloncesto acierta en el 40% de sus lanzamientos. Si en un partido intenta 20 lanzamientos ¿cuál es la probabilidad de que meta exactamente 11 canastas? Para conocer la respuesta hay unos cálculos estándar que pueden realizarse. Hay también otro método que, aunque en este caso es opcional, a veces es el único modo de abordar el problema. En el caso planteado, supondría pedir al jugador que jugara rápidamente unos 10.000 partidos para que pudiéramos determinar el porcentaje de veces en las que mete exactamente 11 canastas.

Aunque es claramente impracticable para un jugador de baloncesto humano, este método, llamado de Montecarlo, puede realizarse fácilmente en un ordenador. Basta con pedir al ordenador que genere un número entero aleatorio comprendido entre 1 y 5, y que mire si dicho número es 1 o 2. Como 2 es el 40% de 5, si sale 1 o 2 lo interpretaremos como la simulación de un acierto del jugador, mientras que si sale 3, 4 o 5 lo interpretaremos como un fallo. Luego pediremos al ordenador que genere 20 de tales números aleatorios entre 1 y 5, y que mire si exactamente 11 de ellos son 1 o 2. Si es así, lo interpretaremos como si el jugador simulado hubiera metido exactamente 11 canastas de 20 intentos cuando su porcentaje de aciertos es del 40%. Por fin, pediremos al ordenador que realice este pequeño ejercicio 10.000 veces y que contabilice el número de veces en que 11 de los 20 intentos del partido simulado se convierten en canasta. Si dividimos este número por 10.000 tendremos una muy buena aproximación de la probabilidad teórica en cuestión.

Para apreciar la utilidad de la simulación es bueno realizar una por uno mismo. (No se preocupe. No hace falta ordenador, basta con una moneda).

Imagine que un gobierno sexista de un cierto país le contrata como asesor. Acaba de adoptar una política que obliga a las parejas a tener hijos hasta que les nazca el primer varón, momento en el que han de cesar de procrear. Lo que quieren saber los gobernantes de ese país es: ¿cuántos hijos tendrá la familia media como resultado de esta política?, y ¿cuál será la distribución de sexos? En vez de hacer una recopilación de datos estadísticos, para lo cual necesariamente se tardarían años, uno puede lanzar una moneda al aire para tener una muestra suficientemente grande que permita hacer una estimación. Interpretando la cara como varón (V) y la cruz como hembra (H), uno lanza la moneda al aire hasta que sale la primera cara y apunta el número de lanzamientos, esto es, el número de hijos de la familia. La sucesión HHV corresponde a dos chicas seguidas de un chico, V corresponde a un hijo único varón, etc. Repítase este procedimiento 100 o 1.000 veces para producir 100 o 1.000 «familias» y calcúlese el número medio de hijos de cada familia y la distribución por sexos. Puede que usted, y también los funcionarios del país, encuentren sorprendente la respuesta.

La aplicación de los métodos de Montecarlo facilita grandemente los estudios de grandes sistemas, las situaciones que se dan en problemas de colas y planificación de horarios, y en fenómenos físicos, tecnológicos y matemáticos. Desde las rebajas de los grandes almacenes a los laboratorios de turbulencia aeronáutica, todo el mundo simula. Generar números aleatorios en un ordenador y manejar luego las simulaciones probabilísticas basadas en ellos es más fácil y barato que tratar con fenómenos aleatorios reales. La única advertencia es que no hay que olvidar que existe una clara diferencia entre el modelo o simulación de un fenómeno y el fenómeno real propiamente dicho. No es lo mismo tener un hijo que lanzar una moneda al aire.

En relación a esto es útil la siguiente representación esquemática de la simulación —y, hasta cierto punto, de la matemática aplicada en general—. El proceso puede dividirse en cinco estadios: el primero es la identificación del fenómeno real que nos interesa; a continuación, la creación de una versión idealizada de dicho fenómeno; en tercer lugar, la construcción de un modelo matemático basado en dicha versión simplificada; luego, la realización de una serie de operaciones matemáticas con el modelo para

obtener predicciones y, por último, la comparación de estas predicciones con el fenómeno original para ver si concuerdan. (Véase la entrada sobre *La filosofía de la matemática*).

Muchas aplicaciones de la matemática son inmediatas, pero lamentablemente es muy fácil, especialmente en las ciencias sociales, confundir el modelo propio con la «realidad» y atribuir a esta última alguna propiedad que sólo existe en el modelo. He aquí un ejemplo simple tomado del álgebra elemental: Jorge puede realizar una tarea en 2 horas y Marta emplea 3 horas en realizar la misma tarea. ¿Cuánto tardarán trabajando los dos a la vez? La respuesta «correcta» de 1 hora y 12 minutos supone que, trabajando juntos, la presencia de uno no estorba ni estimula el trabajo del otro. En este caso, y en muchísimos otros, la certeza de las conclusiones matemáticas derivadas del modelo no siempre es extensiva a las suposiciones, simplificaciones y datos que uno ha empleado en la construcción del modelo. Éstos no se dejan manejar bien, son nebulosos y totalmente falibles, a pesar de las afirmaciones, fastidiosamente autosuficientes a veces, de sociólogos, psicólogos y economistas. Al igual que la señora Brown, la mujer perfectamente ordinaria de Virginia Woolf, la realidad es infinitamente compleja e imposible de capturar por completo en ningún modelo matemático.

[La respuesta al problema de simulación es que el número medio de hijos por familia es dos, un varón y una hembra].

Música, pintura y digitalización

La música y el número van de la mano desde los tiempos de Pitágoras y sus discípulos. Ellos fueron los primeros en darse cuenta de la relación entre la proporción matemática y los sonidos armónicos. Pulsando cuerdas tensas de longitudes proporcionales a los números enteros se producían tonos eufónicos, y prolongando las longitudes de las cuerdas en proporción a los números enteros se producía una escala entera.

Saltándome por pura ignorancia más de dos milenios, anotaré que la obra del matemático de principios del siglo XIX, Joseph Fourier, amplió considerablemente estos conocimientos y proporcionó un formalismo que permite describir matemáticamente cualquier sonido musical como combinación de funciones trigonométricas periódicas. El tono, el volumen y la calidad de un sonido se pueden relacionar, respectivamente, con la frecuencia, la amplitud y la forma de las funciones periódicas que lo representan.

La moderna música electrónica se basa, entre otras cosas, en estas ideas. Desde las composiciones por ordenador de John Cage a los últimos avances en las técnicas de grabación, la matemática está íntimamente ligada a la música y a su tratamiento.

Mucho menos matemática que la música, la pintura presenta pocos paralelismos en este desarrollo. Quizás el uso de la geometría proyectiva y la perspectiva por Alberto Durero, Leonardo da Vinci y otros artistas del Renacimiento podría considerarse en cierto modo análogo a las proporciones de los pitagóricos, y la evolución de la infografía y la geometría fractal podría tener, al menos superficialmente, un cierto parecido con la música digital.

Este último adjetivo, «digital», apunta a la razón por la que la asociación

entre música y matemática se ha dado en una forma más natural que entre pintura y matemática. La música es digital, o se puede poner fácilmente en forma digital o discreta, mientras que con la pintura no se ha podido hacer hasta hace muy poco. Nótese, por ejemplo, que tanto en la notación musical como en la notación numérica, la posición es importante; el lugar de la nota en el pentagrama determina su tono. De hecho, desde la escritura de notas y partituras hasta el empleo de simetrías y repeticiones en obras tan diversas como las de Bach y Cage o hasta la tecnología acústica de la construcción de órganos y discos compactos, los números y la matemática siempre han tenido un papel de apoyo importante en música. Por el contrario, «pintar por números» denota vulgaridad, al igual que la frase «procesamiento de cuadros».

Esta diferencia puede expresarse claramente en términos de ordenadores. Las máquinas primitivas en las que la salida variaba gradualmente con una magnitud física continua como el voltaje o la presión se denominaron ordenadores analógicos. Esto contrasta con los hoy más conocidos ordenadores digitales, cuya salida depende completamente de si una cierta condición lógico-electrónica se satisface o no. Sin embargo, como sugieren también las anteriores observaciones esquemáticas sobre la música y la pintura, la distinción se viene abajo si se abusa de ella. Un parche de tambor vibrando se asocia más naturalmente a un instrumento analógico que a uno digital, y la edición de gráficos basada en técnicas de píxel a píxel es en sí un procedimiento digital. (Un píxel es el menor punto de luz visible en una pantalla; la mayoría, al menos hace unos años, tienen unos 200×300 o 60.000 píxels).

Sin embargo, es útil resaltar este contraste, y algunos de sus aspectos son manifiestos en bastantes contextos. Los velocímetros, por ejemplo, pueden ser digitales y dar una lectura numérica, o analógicos y dar como lectura un rectángulo que se alarga (o la rotación de una aguja). La salida analógica es menos precisa pero da la velocidad enmarcada en una cierta perspectiva. Un 82 y un 28 son casi indistinguibles, pero en cambio un rectángulo largo es muy distinto de uno corto, y si la longitud está cambiando también esto es manifiesto. Algo similar ocurre con los relojes; los digitales son más precisos, pero carecen de todas las asociaciones que comporta una esfera clásica.

He observado una diferencia parecida en el impacto que me produce un texto dependiendo de si lo leo directamente en el monitor de un ordenador o si leo una copia en papel sintiendo el crujir de las hojas. Como antes, este segundo caso, más parecido a una presentación analógica, crea un ambiente más sugerente que la correspondiente presentación digital en la pantalla. Consigo mejor hacerme una idea de la estructura, el peso y las proporciones de la obra si puedo tocar el manuscrito. Análogamente, aunque esté absolutamente a favor de la utilización creciente de las calculadoras digitales en clase, pienso que sus resultados numéricos se han de cotejar con el sentido común y que hay que desarrollar la capacidad de estimar y comparar magnitudes.

La digitalización de dispositivos que en su día fueron analógicos comporta a menudo un aura de artificiosidad. El sonido de instrumentos y voces se puede modificar fácilmente, o incluso ser sintetizado por los ingenieros y mezcladores de sonido, cuyo producto, el disco compacto, puede contener una música que nunca antes se ha oído en un entorno natural. Más sorprendente aún es el hecho de que las fotografías, que se han considerado siempre como pruebas documentales de la realidad, puedan alterarse por técnicas digitales parecidas, aunque más sofisticadas. Una fotografía de un hombre y una mujer puestos contra una pared de ladrillos, por ejemplo, ya no significa que ese hombre y esa mujer hayan estado alguna vez juntos contra esa pared. Las fotografías pueden descomponerse y reconstruirse electrónicamente, de modo que los colores, tejidos, superficies, vestidos y caras pueden modificarse a voluntad, siendo el resultado indistinguible de una fotografía «real».

Aunque no quiero poner demasiado énfasis en la distinción un tanto nebulosa entre digital y analógico, ésta tiene también su importancia para el modo cognitivo. Persiste en el mundo de la informática, por ejemplo, en la diferencia en la entrada de datos por medio del teclado (que es la que prefiero) o por medio de un ratón. Esto último tendría un aire más analógico. En matemática normalmente prefiero también la información expresada digitalmente en términos de palabras y símbolos a las expresiones analógicas que contienen figuras y diagramas. A pesar de todo, independientemente de las preferencias personales, la distancia entre los modos matemáticos de

comunicación y los puramente verbales es menor de lo que le parece a la mayoría de la gente, y por ello me sorprende que los libros de matemáticas raramente exploten la prodigiosa facilidad que tenemos todos con el lenguaje, tanto los numéricos como los anuméricos. Si casi siempre las ideas matemáticas se pueden expresar con palabras ¿por qué no se usa más a fondo la más general y potente de las herramientas?

El abuso de acompañamientos gráficos distrae la atención. Tengo un amigo que se ha quedado tan prendado de sus programas de gráficos por ordenador que es incapaz de escribir una carta personal sin llenarla de ilustraciones y polígonos de frecuencias. (Admito que lo mismo se hace con las palabras, a veces en forma de fastidiosos apartes entre paréntesis como éste). Además de distraer la atención del tema que se está tratando, tales figuras geométricas también intimidan a muchas personas con sensibilidad literaria y/o un pobre bagaje matemático. No hace falta decir que estos prejuicios explican la relativa escasez de ilustraciones de este libro.

Me detendré aquí, pues empiezo a parecerme peligrosamente a esas personas que lo explican todo, desde la deuda nacional a la poesía épica de Islandia, murmurando cualquier trivialidad acerca de los hemisferios cerebrales derecho e izquierdo. Baste decir que la música, la pintura y la matemática pueden emplear distintas mezclas de técnicas y talentos, pero que para las tres hacen falta cerebros enteros.

Notación

Hay seis cosas a tener en cuenta para elegir una buena notación matemática, y me referiré a ellas como 1, b, III, cuatro, E y vi, respectivamente. La lista anterior, adaptada de un chiste del cómico George Carlin, es un ejemplo de notación mal elegida. Sin embargo, excepto en casos extremos como éste, la mayoría de la gente cree, si es que ha pensado alguna vez sobre ello, que la notación matemática no puede inducir a confusión ni tampoco aportar nuevas ideas, sino que se trata simplemente de un aspecto superficial, casi cosmético, de la matemática. Cómo denotar líneas, ángulos, puntos, números y otros objetos matemáticos: ¡vaya cosa más banal e irrelevante! La eterna preocupación por la notación del pedante no hace sino reforzar esta creencia natural que, sin embargo, no siempre tiene un fundamento sólido. Los sistemas de notación son a veces algo más que tipologías grandilocuentes y convenios vacíos.

De hecho, la invención de una notación conveniente y flexible es, a menudo, más fructífera que una demostración del más profundo de los teoremas. Se podría incluso escribir una especie de esbozo de historia de las matemáticas dedicada por completo a la introducción de nuevas notaciones importantes. Los números romanos (y sus chapuceros primos griegos), a pesar de su persistente y pretencioso uso en contextos como Super Bowl XXIV y Copyright MCMLXXXIX, no se prestan tan bien al cálculo como los números arábigos que les sustituyeron. (Véase la entrada sobre *Números arábigos*). La discusión acerca de números y magnitudes concretas no se puede generalizar fácilmente sin la introducción de una notación para las variables (inventada por François Viète). Y, análogamente, es difícil manipular simbólicamente figuras geométricas sin las herramientas de

notación de la geometría analítica (René Descartes y Pierre Fermat).

Verdaderamente no había ningún teorema formal relacionado con la introducción de los simbolismos anteriores, pero es absolutamente cierto que no fueron simple cosmética. Su introducción y posterior adopción universal codificó las ideas e intuiciones de unas personas muy listas y eruditas y las hizo accesibles a todo el mundo. En unos pocos minutos, los estudiantes de cuarto de básica (¿habría que decir, quizá, de IV de básica?) pueden realizar unos cálculos con los que los profesores de la Europa medieval tardaban horas. Utilizando variables, los estudiantes de séptimo de básica pueden plantear y resolver ecuaciones algebraicas cuyas soluciones eran en algunos casos desconocidas tanto para los matemáticos de la antigua Grecia como para los mismos profesores medievales. Mediante las técnicas de la geometría analítica, un estudiante de primer año de universidad puede formarse ideas y obtener propiedades de las figuras geométricas que escapaban a quienes tenían que limitarse a los medios de expresión clásicos.

En la misma vena, la notación de la teoría de conjuntos, común en todas las ramas de la matemática, es lo bastante simple para ser enseñada en la enseñanza básica. A pesar de esto, el capítulo introductorio de un sinfín de textos de cualquier nivel contiene unas diez o doce páginas que vienen a decir más o menos lo mismo que las frases siguientes. Escribimos $p \in F$ para indicar que p es un elemento del conjunto F y $F \subset G$ quiere decir que todo elemento de F es también un elemento de G . Dados dos conjuntos A y B , $A \cap B$ es el conjunto de los elementos que pertenecen a la vez a A y a B , $A \cup B$ es el conjunto de los elementos que pertenecen a A , a B o a ambos, y A' es el conjunto de los elementos que no pertenecen a A . El conjunto vacío, conjunto que no tiene ningún elemento, se indica unas veces por $\emptyset$ y otras por $\{ \}$, unas llaves que no encierran nada. Fin del cursillo.

Las notaciones anteriores, así como las desarrolladas desde entonces — para las derivadas, integrales y series en el cálculo, las notaciones de operadores y matrices en varias disciplinas matemáticas, los símbolos de los cuantificadores y los conectivos lógicos, la teoría de categorías, que en cierto sentido no es más que una teoría de la notación—, tienen todas ellas tres propiedades (A, 2 y III) que explican su adopción y uso. Son sugerentes, la relación entre los símbolos se corresponde de una manera natural con la

relación entre los objetos matemáticos (Y^5 , por ejemplo, significa $Y \times Y \times Y \times Y \times Y$). Generalmente, también, los buenos simbolismos son fácilmente manejables y emplean ciertas reglas o algoritmos (ya sea para multiplicar números o para resolver sistemas lineales de ecuaciones diferenciales). Por último, las notaciones eficaces abrevian, contienen mucha información crítica en una forma compacta sin elementos extraños ni redundantes.

Naturalmente, la notación como sistema formalizado de representación trasciende la matemática. La contabilidad por partida doble (en la que cada transacción se entra a la vez como débito y como crédito) transformó la práctica de la teneduría de libros, los símbolos químicos simplificaron muchísimo la expresión de las reacciones químicas, y los diagramas de Feynman para las interacciones subatómicas clarificaron la mecánica cuántica. Hay una infinidad de otros ejemplos de notaciones, de los cuales el más omnipresente, de alcance universal, es el alfabeto y la lengua escrita que disponemos gracias a él.

La notación científica

Rápidamente, ¿qué es mayor, 381000000000 o 98200000000? Y ¿qué es menor, 0,00000000034 o 0,0000000085? Sin contar los ceros, o sin la presencia de los espacios separadores, hace falta algo más que una simple mirada para captar la magnitud de 9.450.000.000.000.000, el número de metros de un año luz, o la pequeñez de 0,000 006 5, la longitud de onda de la luz roja expresada en metros. La notación científica es un convenio que nos permite escribir estos números, muy grandes o muy pequeños, de modo que sus magnitudes relativas sean manifiestas. Además, nos ahorra el esfuerzo de tenerlos que pronunciar.

Para comprender esta notación hay que recordar que 10^1 es otra manera de escribir 10, que 10^2 es 100, que 10^6 es 1.000.000 y que en general 10^N es la unidad seguida de N ceros. De ahí pues que 4×10^3 es 4×1.000 o 4.000, y que $5,6 \times 10^6$ es $5,6 \times 1.000.000$ o 5.600.000. Para expresar números pequeños se usan los exponentes negativos, y definimos 10^{-1} como 0,1, 10^{-2} como 0,01, 10^{-6} como 0,000 001, y en general 10^{-N} como la unidad precedida de un 0, una coma decimal y N – 1 ceros más (o también $10^{-N} = 1/10^N$). De ahí que $5,7 \times 10^{-3}$ sea $5,7 \times 0,001$ o 0,0057, y $9,1 \times 10^{-8}$ sea $9,1 \times 0,000\ 000\ 01$ o 0,000 000 091.

Es parte del convenio que a la izquierda de la coma sólo haya un dígito distinto de 0. Así 48.700 se escribe $4,87 \times 10^4$ y no $48,7 \times 10^3$, y 0,000 000 23 se escribe $2,3 \times 10^{-7}$ y no 23×10^{-8} , aunque en cada caso las dos expresiones indican el mismo número. Análogamente, se escribe 239.000.000

como $2,39 \times 10^8$; 59.700.000.000.000 como $5,97 \times 10^{13}$; 0,000 031 como $3,1 \times 10^{-5}$; y 0,000 000 002 5 como $2,5 \times 10^{-9}$.

Además de simplificar el reconocimiento, la notación científica hace más fáciles los cálculos aproximados de orden de magnitud, pues nos permite usar la ley de los exponentes, según la cual $10^M \times 10^N$ es igual a 10^{M+N} o, en particular, que $10^5 \times 10^8 = 10^{13}$. Por ejemplo, como el mundo tiene aproximadamente 5×10^9 habitantes y cada persona tiene en la cabeza una media de aproximadamente $1,5 \times 10^5$ cabellos, en la Tierra hay aproximadamente $(5 \times 10^9) \times (1,5 \times 10^5)$, o $7,5 \times 10^{14}$ cabellos humanos. (Me apartaré un poco del tema para observar que la ley de la suma de exponentes explica el uso de los exponentes fraccionarios. Si $X^? \times X^? = X^1$, entonces $X^?$ ha de ser igual a $\sqrt{X}$, pues al elevarlo al cuadrado da X^1 . Pero $X^{1/2} \times X^{1/2}$ también da X^1 , por aplicación formal de la ley de suma de los exponentes. Así pues, $X^{1/2}$ se define como $\sqrt{X}$).

Para poder traducir al español la peluda cifra de los $7,5 \times 10^{14}$ y otras parecidas, hemos de echar mano de una especie de equivalente oral de la notación científica, esto es, de algunos vocablos numéricos tradicionales. Mil es 10^3 ; un millón es 10^6 ; un billón 10^{12} ; y un trillón, 10^{18} . Para hacerse una idea visceral de las diferencias entre estos números puede servir pensar que mil segundos son aproximadamente 17 minutos, un millón de segundos tardan unos 11 1/2 días en pasar, mil millones de segundos son unos 32 años y un billón de segundos, casi 32.000 años.

Si se conocen estas equivalencias y unos cuantos datos comunes (la población de Estados Unidos, la del mundo, la distancia entre las costas este y oeste) es más fácil evaluar la importancia y la exactitud de algunos números grandes y valorar de un modo racional la magnitud de los riesgos. Para poner un ejemplo tópico, el número de servicios sexuales realizados anualmente por las prostitutas norteamericanas se estima en unos trescientos millones (o si lo prefieren 3×10^8). ¿Es una cifra razonable? Y, si no lo es, ¿es mayor o menor el verdadero número? Además de a los académicos y los adolescentes, la respuesta interesa también a los epidemiólogos del SIDA que intentan

explicar la aparente escasez de hombres cuya infección pueda ser atribuida a tales contactos. O considérese también el medio billón (5×10^{11}) de dólares (según algunas estimaciones) del escándalo de las cajas de ahorros y empréstitos ¿Duda alguien de que permaneció invisible durante tanto tiempo porque era un «asunto de números» y no un «asunto de personas»?

A menudo se emplean prefijos en lugar de exponentes. Los más corrientes son kilo- para 10^3 , mega- para 10^6 , giga- para 10^9 y tera- para 10^{12} . En sentido descendente el prefijo mili- indica 10^{-3} ; micro-, 10^{-6} ; nano-, 10^{-9} ; y pico-, 10^{-12} . Un nanosegundo es una milmillonésima de segundo y guarda con el segundo la misma proporción que éste con los treinta y dos años del gigasegundo.

El caprichoso vocablo «googol», acuñado por el matemático Edward Krasner para indicar la unidad seguida de cien ceros (10^{100}), nos lleva a reinos aún más insondables. Un googolplex es incomparablemente más inimaginable, pues se define como la unidad seguida de un googol de ceros: $10^{(\text{googol})}$. Es difícil encontrarse con conjuntos de objetos reales que tengan tantos elementos. El físico Arthur Eddington escribió en cierta ocasión que en todo el universo había aproximadamente $2,4 \times 10^{79}$ partículas (menos de un googol) y, aunque la física en la que se basaba su estimación está muy superada, sigue siendo válido que, mientras no abandonemos las cosas reales, estos grandes números nos bastan y nos sobran. No obstante, si nos ponemos a contar *posibilidades*, es fácil superar incluso estos números. Si lanzamos una moneda al aire sólo 1.000 veces, por ejemplo, el conjunto de sucesiones posibles de caras y cruces es, por la regla del producto, $2^{1.000}$, o aproximadamente 10^{300} , que es ¡el cubo de un googol!

[Aunque la notación científica es útil en muchos contextos, no lo es, por supuesto, en muchos otros. En una guardería especial para niños «superdotados» de Filadelfia vi en cierta ocasión cómo los académicamente obsesos padres de un genio en pañales abrían los brazos y le decían, «Te queremos 10^8 ». Vestidos al estilo *hippy* de los sesenta, sólo estaban haciendo una payasada, pero yo no podía dejar de preguntarme si más tarde, al llegar a casa, no se pondrían a escuchar la vieja canción de folk «Viviré contigo

$3,65 \times 10^2$ días» o le contarían a su hijo cuentos de «Las 10^3 y 1 Noches»].

Números arábigos

Un mercader alemán del siglo XV preguntó a un eminente profesor dónde tenía que mandar a su hijo para que recibiera una buena formación mercantil. El profesor le contestó que las universidades alemanas bastarían para que el muchacho aprendiera a sumar y a restar, pero que para aprender a multiplicar y dividir tenía que ir a Italia. Antes de sonreír indulgentemente, intentad multiplicar, o aunque sólo sea sumar, los números romanos CCLXIV, MDCCCIX, DCL y MLXXXI, sin convertirlos previamente.

Puede que los números sean eternos e invariantes, pero los numerales, o símbolos empleados para representarlos, no lo son y la anécdota anterior sirve para ilustrar qué fácil es dar por supuesto el sistema de numeración indoarábigo que usamos actualmente. La historia de los sistemas de numeración es muy larga y va desde la prehistoria hasta la adopción de nuestro actual sistema, en el Renacimiento. Los protagonistas de la historia son escribas, contables, monjes y astrónomos anónimos que descubrieron los principios de la representación sistemática de los números.

Estos principios —simbolización abstracta (en contraposición a representaciones concretas con guijarros, por ejemplo), notación posicional (826 es muy distinto de 628 o 682), un sistema de base multiplicativa [el numeral 3.243 en nuestro sistema de base 10, por ejemplo, se interpreta como $(3 \times 10^3) + (2 \times 10^2) + (4 \times 10) + (3 \times 1)$, a diferencia de su interpretación en un sistema de base 5, donde valdría 428: $(3 \times 5^3) + (2 \times 5^2) + (4 \times 5) + (3 \times 1)$], y el santo grial, es decir, el cero (que permite distinguir fácilmente entre 36, 306, 360 y 3.006)— son una parte esencial, aunque casi invisible,

de nuestra herencia cultural.

[Una «aplicación» no convencional de la expresión de los números en una base distinta se tiene cuando alguien ha de celebrar un cumpleaños no deseado, como el 40° por ejemplo. En el sistema de base 12, 40 se escribe 34: $(3 \times 12) + (4 \times 1)$. En una base distinta el cumpleaños pierde parte de su importancia artificial. Observaciones parecidas valen para la epidemia de necesidad numerológica que espero provocará la proximidad del cambio de milenio en el 2000, o el 2001 según los puristas. Sin embargo, cambiar de base no sirve siempre. La insensatez que rodea al 666 se habría prendido de algún otro número en un mundo de base 5, donde 666 tendría la expresión inocua 10.131: $(1 \times 5^4) + (0 \times 5^3) + (1 \times 5^2) + (3 \times 5^1) + (1 \times 1)$. Aplicando el principio de conservación de la superstición, 444 —ciento veinticuatro en nuestro sistema decimal— podría haber tenido importancia mística en ese mundo].

Para expresar los números se han empleado muchas maneras concretas: guijarros de distintos tamaños, cuerdecitas anudadas de distintos colores, el omnipresente ábaco y, finalmente, el más personal de los ordenadores personales: las manos y los pies. Contar con los dedos de las manos o con los de las manos y los pies (¿dígitos de quita y pon?) es una práctica casi universal que a la larga dio lugar a las bases escritas más comunes. Nuestro sistema de base 10 es, en efecto, consecuencia de esto, mientras que las palabras que en francés significan 20, 80 y 90 —*vingt*, *quatre-vingts* y *quatre-vingt-dix*— sugieren un sistema de base 20 más antiguo. Hace 1500 años los mayas, uno de los cuatro pueblos que inventaron el principio de la notación posicional, usaban también un sistema de base 20 y crearon un calendario más preciso que el gregoriano que usamos en la actualidad. Incluso el antiguo sistema de base 60 de los sumerios y babilonios, que ha sobrevivido en nuestro modo de medir el tiempo, los ángulos y las posiciones geográficas, se derivó probablemente del contar con los dedos.

Siguiendo con nuestro esbozo, observamos que hace aproximadamente 2.000 años los chinos inventaron un sistema de numeración posicional escrito basado en las potencias de 10. Unos 500 años después la gente del sur de la India llegó independientemente al mismo descubrimiento, y pronto fueron más allá inventando el cero, un símbolo que revolucionó el arte de

representar y manipular los números. Antes de llamarse así, el cero se solía representar por un espacio vacío en un numeral o en un ábaco. Después, el símbolo correspondiente indicaba que había un espacio o, lo que es lo mismo, que faltaba algo. Finalmente, quedó claro que los números se definían por sus propiedades y que el cero tenía tantas como cualquier otro.

Los chinos tomaron prestado de los indios el concepto de cero, igual que los árabes, que, al cabo del tiempo, exportaron todo el sistema a la Europa occidental. El sistema de numeración indo-arábigo se puede colocar con todo merecimiento entre los mayores descubrimientos técnicos de la humanidad, junto con la invención de la rueda, el fuego y la agricultura.

[A propósito, la suma de los números citados en el primer párrafo es MMMDCCCIV. ¿Cuál es su producto?].

Números y códigos binarios

Los números binarios consisten en una sucesión de unos y ceros, y aunque ya eran conocidos por los matemáticos de la antigua China, fueron estudiados en serio por vez primera por el matemático y filósofo alemán Gottfried Wilhelm von Leibniz, motivado por consideraciones metafísicas sobre el ser y el no ser. Como muchos fenómenos (¿todos?) pueden reducirse a sucesiones complejas de dicotomías abierto-cerrado, encendido-apagado o sí-no, y como al menos los ordenadores funcionan así, hace ya mucho tiempo que los números y códigos binarios han descendido del reino de la metafísica al de lo mundano.

En cualquier caso, ¿cómo haremos para acomodar nuestros números arábigos a este ropaje más austero de 0 y 1? Aquí servirá más un ejemplo que una explicación. El número 53 se puede expresar como $32 + 16 + 4 + 1$, donde cada término es una potencia de 2 (se entiende que el 1 es la potencia cero de 2, 2^0). Tomaremos 110101 como representación binaria de 53, indicando cada 1 o 0 la presencia o la ausencia de la correspondiente potencia de 2. Esto es, $53 = (1 \times 2^5) + (1 \times 2^4) + (0 \times 2^3) + (1 \times 2^2) + (0 \times 2^1) + (1 \times 2^0)$. Como ocurre con el sistema de numeración arábigo, la posición de un dígito determina su valor.

Más ejemplos: el número $83 = 64 + 16 + 2 + 1$, por lo que expresado en el sistema binario es 1010011: $(1 \times 2^6) + (0 \times 2^5) + (1 \times 2^4) + (0 \times 2^3) + (0 \times 2^2) + (1 \times 2^1) + (1 \times 2^0)$. Análogamente, $217 = 11011001$. Los números del 1 al 16 expresados en notación binaria son: 1, 10, 11, 100, 101, 110, 111, 1000, 1001, 1010, 1011, 1100, 1101, 1110, 1111, 10000. Y actuando en sentido inverso, 110100 se convierte en 52, 1111100 en 124, y 1000000000

en 512 (2^9). Una vez sabemos expresar los números arábigos en forma binaria, podemos aplicar las mismas reglas y algoritmos aritméticos para trabajar con ellos, con la única salvedad de que hemos de recordar que tratamos con potencias de 2 en vez de potencias de 10.

Sin limitarnos sólo a los números, los códigos binarios pueden usarse también con más generalidad en modos diversos. Por ejemplo, si asignamos valores de verdad (1 para verdadero, 0 para falso) a los enunciados, entonces las operaciones elementales de la lógica —negar un enunciado, unir dos enunciados con un «y», un «o» o un «si..., entonces...», etc.— se pueden realizar fácilmente mediante operaciones sintácticas sencillas o, físicamente, mediante circuitos electrónicos sencillos. (Véase la entrada sobre *Tautologías*). A un enunciado precedido de un «no» se le asigna el valor de verdad de 0 o 1 según sea el valor de verdad del enunciado original 1 o 0. El enunciado formado por unión de otros dos enunciados por medio de un «y» tiene un valor de verdad de 1 sólo si los dos enunciados más cortos tienen el valor de verdad 1. [Estas operaciones lógicas se llaman booleanas en honor del matemático inglés del siglo XIX, George Boole, quien, según la apreciación hiperbólica (¿hiperboolica?) de Bertrand Russell, descubrió la matemática pura].

Codificar las letras y otros símbolos como secuencias de ceros y unos no presenta ningún problema, ya que a cada carácter se le asigna una sucesión distinta de dígitos binarios o bits. Según el convenio estándar ASCII para ordenadores, cada símbolo tiene un código de 8 bits (8 bits = 1 byte), y en total hay 256 (2^8) códigos de éstos —uno para cada una de las 52 letras, mayúsculas y minúsculas, para los números del 0 al 9 y para los signos de puntuación, símbolos aritméticos y de control, y otros—. El código de la P es 01010000; el de la V, 01010110; el de la b, 01100010; la t es 01110100; el de «, 00100010; el de &, 00100110, etc. Estos códigos se usan en el procesamiento de textos, en cuyas aplicaciones los símbolos no se suelen operar en el sentido aritmético del término, sino que sólo se presentan ordenados como texto.

Para ilustrar cómo un solo número binario puede codificar un gran conjunto de información, piénsese que esta entrada contiene algo más de 5.600 símbolos —entre letras, dígitos, espacios blancos y signos de

puntuación—. Cada uno de ellos tiene un código de 8 bits, con lo que si los encadenamos todos, resulta una sucesión de aproximadamente 45.000 bits que podemos tomar como representación binaria del artículo. Podemos hacer lo mismo con todo el libro y obtendremos su código binario. O siendo aún más ambiciosos, si ordenamos todos los libros de la Biblioteca Nacional, alfabéticamente según el nombre del autor y cronológicamente por la fecha de publicación, y encadenamos sus sucesiones, obtendremos el número binario que representa toda la información de la Biblioteca Nacional.

La teoría matemática de la información estudia la cantidad de información que representa una cadena binaria (o al menos da una definición muy útil de dicha cantidad). Esa teoría es un fértil campo de investigación con muchas aplicaciones a la biología, la lingüística y la electrónica, y se expresa en términos de bits, de modo que cada bit de información comporta una elección binaria. [Por tanto 5 bits implican, por ejemplo, 5 de tales elecciones y bastan para distinguir entre 32 (2^5) alternativas, pues hay 32 (2^5) posibles sucesiones de síes y noes de longitud 5]. Los bits sirven también de unidades de medida numérica de conceptos tales como la entropía de las fuentes de información, la capacidad de los canales de comunicación y la redundancia de los mensajes.

Desde la teoría de la información a los puntos y rayas del código Morse y a las líneas de distinto grosor de los códigos de barras de los supermercados, los números y los códigos binarios están hoy por todas partes. (El mundo está entrando en los bits y los ordenadores personales). Sin embargo, pienso que se comprenden mejor cuando se les considera con un poco del aprecio original de Leibniz por su primacía metafísica. La información, los ordenadores, la entropía, la complejidad —todos estos conceptos e ideas fundamentales— proceden en parte del código más elemental de todos, 1 o 0, sustancia y nada, yin y yang, ser o no ser. Es fácil verse arrastrado por una orgía de oposiciones sinónimas, así pues me detendré aquí, pero acabaré señalando que un universo que fuera todo sustancia sería indistinguible de otro que estuviera completamente vacío, con lo que una cierta dicotomía binaria es una condición previa para un universo no trivial y también para el propio pensamiento.

Números imaginarios y números negativos

Se podría esbozar una historia abreviada de los distintos conjuntos de números considerando las soluciones de diferentes tipos de ecuaciones algebraicas. (Véase también la entrada sobre *Números arábigos*). La ecuación $2X + 5 = 17$ tiene por solución el número entero positivo 6. Sin problemas por el momento. Pero si queremos resolver la ecuación $3X + 11 = 5$ hemos de ir más allá de estos números tan simples. La solución de la ecuación anterior es un número negativo, -2 , aunque se tardó bastante tiempo en dar el paso decisivo y llamar número a -2 . Los matemáticos tardaron mucho en entender la categoría de los números negativos y, aún hoy, sus propiedades son un poco misteriosas para quienes se inician en el álgebra. Nadie tiene demasiados problemas con los números negativos en sí. Quince grados bajo cero es perfectamente comprensible, tanto visceral como intelectualmente. Pero ¿por qué un número negativo por un número negativo da un número positivo?

La respuesta es formal. El producto de dos números negativos se define positivo para que estos números obedezcan las mismas leyes aritméticas que los enteros positivos. Los ejemplos financieros son muy ilustrativos. Suponga que le pagan semanalmente 10.000 ptas. de una pequeña pensión, que usted guarda puntualmente bajo el colchón. Así dentro de 7 semanas tendrá bajo el colchón 70.000 ptas. más que hoy ($7 \times 10.000 \text{ ptas.} = 70.000 \text{ ptas.}$), mientras que hace 5 semanas había 50.000 ptas. menos que hoy ($-5 \times 10.000 \text{ ptas.} = -50.000 \text{ ptas.}$). Suponga ahora que, pasados algunos años, su pensión se ha acabado y que tiene que pagar 10.000 ptas. a la semana que saca de su escondrijo bajo el colchón. Entonces, al cabo de 8 semanas habrá en el escondite 80.000 ptas. menos que hoy ($8 \times -10.000 \text{ ptas.} = -80.000 \text{ ptas.}$)

mientras que hace 3 semanas había 30000 ptas. más que hoy (-3×-10.000 ptas. = 30.000 ptas.).

Así pues, completamos nuestro sistema de números enteros positivos con los enteros negativos. Pero su conjunto, los enteros, todavía no basta para resolver ecuaciones como $5X - 1 = 7$, cuya solución, $8/5$, es una fracción (un número racional). Igual que hicimos antes, somos hospitalarios e incorporamos todas las fracciones a nuestro sistema numérico. Pero incluso los números racionales son insuficientes para calmar nuestra sed de soluciones. La ecuación $X^2 - 2 = 0$, por ejemplo, tiene como solución la raíz cuadrada de 2, que no es un número racional. Como tampoco lo es la solución de $4X^3 - 7X + 11 = 0$. Quizá si incorporáramos a nuestro sistema numérico todos estos números algebraicos y también todos los números irracionales podríamos llegar a resolver todas las ecuaciones algebraicas.

Falso. Una ecuación tan simple como $X^2 + 1 = 0$ no tiene solución. No hay ningún número real (ni racional ni irracional) tal que $X^2 = -1$ porque el cuadrado de cualquier número real es mayor o igual que 0. ¿Qué podemos hacer? Inventamos un nuevo símbolo, i , que definimos simplemente (al igual que ya hicieron Euler, D'Alembert y otros) como $\sqrt{-1}$, la raíz cuadrada de -1 . Así $i^2 = -1$, y tenemos una solución para nuestra ecuación. Originariamente la letra i se escogió para indicar la naturaleza imaginaria de este número, pero a medida que fue en aumento el grado de abstracción de la matemática, resultó no ser más imaginario que muchas otras construcciones matemáticas. Ciertamente, no sirven para medir cantidades, pero obedecen a las mismas reglas aritméticas que los números reales y, aunque sea sorprendente, permiten formular algunas leyes físicas del modo más natural.

Los números de la forma $a + bi$, donde a y b son reales, constituyen los números complejos, un conjunto de números que incluye a los reales como subconjunto. (Los números reales no son más que aquellos números complejos que tienen $b = 0$. Así los números reales 7,15 y π se pueden escribir, respectivamente, como $7,15 + 0i$ y $\pi + 0i$. El número i se puede escribir como $0 + 1i$). La suma de dos números complejos, pongamos $3 + 5i$ y $6 - 2i$, se define como $9 + 3i$. La resta se define de un modo análogo, y la multiplicación y la división tienen en cuenta el hecho de que $i^2 = -1$. Una

vez construido este sistema numérico ampliado, podemos demostrar el teorema fundamental del álgebra según el cual, aparte de la ecuación ya citada $X^2 + 1 = 0$, cualquier otra ecuación algebraica tiene soluciones en el conjunto de los complejos. (Véase la entrada sobre *La fórmula de la ecuación de segundo grado*). Las ecuaciones $2X^7 - 5X^4 + 19X^2 - 11 = 0$, $X^{17} - 12X^5 + 8X^3 = 0$, y $3X^8 - 26X - 119 = 0$ tienen todas ellas solución en el conjunto de los números complejos. Además, las ecuaciones cuadráticas (aquellas en que la variable está elevada al cuadrado) tienen dos raíces, las ecuaciones cúbicas tienen tres, las cuárticas tienen cuatro y, en general, las ecuaciones polinómicas de grado N (en las que la variable está elevada a la N -ésima potencia como máximo) tienen N raíces.

Aunque los primeros en idear los números imaginarios los usaban sólo formalmente y apenas comprendían lo que estaban haciendo, pronto otros generalizaron las definiciones de las funciones trigonométricas y exponenciales al dominio de los complejos y extendieron el análisis matemático (el cálculo, las ecuaciones diferenciales y otros temas afines) adaptándolo a estas generalizaciones. En particular, se asignó un sentido a la operación de elevar un número a un exponente imaginario. El resultado es una de las fórmulas matemáticas más notables: $e^{\pi i} = -1$, donde e es la base de los logaritmos naturales. Si escribimos la expresión anterior como $e^{\pi i} + 1 = 0$, tenemos las cinco constantes más importantes de la matemática en una sola ecuación. (Ya sé que ocurre lo mismo en $e^{\pi i} = 1$, pero como la potencia 0 de cualquier cosa es 1, la presencia de e , i y π en este caso es superflua). Estos progresos técnicos, entre los que se cuenta la interpretación geométrica de varias operaciones entre números complejos, prepararon el camino para su indispensable uso posterior en la teoría del electromagnetismo y en otras ciencias físicas. Sus avances estimularon también el desarrollo del álgebra abstracta y, en particular, el análisis vectorial y los cuaterniones.

El número i es una prueba de la magnitud de los progresos verdaderos que pueden darse como consecuencia de postular entidades imaginarias. Los teólogos, que han construido elaborados sistemas sobre analogías mucho más débiles, quizá deberían animarse con ello.

Los números primos

Antes de que los progresos de la física nuclear revelaran que el átomo era una sociedad balcanizada de partículas subatómicas, se solía comparar metafóricamente a los números primos con los átomos. Hechos de un material más tenaz (o mejor de un no-material más tenaz) que los átomos físicos, los números primos comparten con ellos su eterna indivisibilidad. Se distinguen de los números compuestos en que éstos pueden expresarse como producto de dos números más pequeños, mientras que con los números primos no se puede. Los números 8, 54 y 323 son compuestos pues son iguales, respectivamente, a 2×4 , 6×9 y 17×19 , mientras que los números 7, 23 y 151 son primos porque no pueden ser descompuestos o factorizados. La primera docena de números primos es 2 (el único primo par: ¿por qué?), 3, 5, 7, 11, 13, 17, 19, 23, 29, 31, 37. Puede demostrarse que sólo hay una manera de expresar un número entero como producto de factores primos. El número 60, por ejemplo, es igual a $2 \times 2 \times 3 \times 5$; el 1.421.625 es igual a $3 \times 5 \times 5 \times 5 \times 17 \times 223$, y el 101, como es primo, simplemente es igual a sí mismo.

Debido a su simplicidad y a su omnipresencia casi tangible, los números primos han fascinado a los hombres desde antes de la antigua Grecia. Una pregunta que surge de manera natural es: ¿cuántos números primos hay? Si continuáramos la lista de números primos empezada más arriba nos daríamos cuenta de que, a medida que fuéramos buscando, los primos se dispersarían cada vez más. Hay más primos entre 1 y 100 que entre 101 y 200. Cabría suponer, pues, que hay un número primo máximo, al igual que hay un elemento con número atómico máximo. Euclides demostró, no obstante, que no hay un número primo máximo y que son, por tanto, infinitos.

La demostración que dio Euclides de esta propiedad es un ejemplo tan bello de lo que suele llamarse demostración indirecta que, arriesgándome a despertar la ansiedad matemática del lector, la reproduciré aquí. Supondremos de entrada que sólo hay un número finito de primos; veremos cómo esta suposición nos lleva a una contradicción. Así pues, escribiremos todos los números primos $2, 3, 5, \dots, 151, \dots P$, donde P representa el mayor número primo. Denotemos por N el resultado del producto de todos ellos, esto es, $N = 2 \times 3 \times 5 \times \dots \times 151 \times \dots \times P$.

Consideremos el número $(N + 1)$ y veamos si es divisible exactamente por 2 (sin resto). Es claro que N es divisible exactamente por 2, pues éste es un factor de N . Por tanto, 2 no divide exactamente a $(N + 1)$, pues queda un resto de 1. También es claro que N es exactamente divisible por 3, pues éste también es uno de sus factores. En consecuencia, 3 no divide exactamente a $(N + 1)$, pues también queda un resto de 1. Lo mismo vale para 5, 7, y todos los números primos hasta llegar a P . Todos ellos dividen exactamente a N y por tanto dejan un resto de 1 cuando el dividendo es $(N + 1)$.

¿Qué significa esto? Como ningún número primo $2, 3, 5, \dots, P$ divide exactamente a $(N + 1)$, o bien $(N + 1)$ es también un número primo, que será mayor que P , o bien es divisible por algún número primo mayor que P . Como hemos supuesto que P era el mayor número primo hemos llegado a una contradicción: la existencia de un número primo mayor que el máximo número primo. Por tanto, la suposición de partida de que sólo hay un número finito de primos ha de ser falsa. Fin de la demostración. QED.

Una proposición más difícil, el teorema del número primo, nos dice aproximadamente con qué frecuencia aparecen los números primos entre los enteros. Si $P(N)$ es el número de primos menores o iguales que N [notemos que $P(10)$ vale 4 pues hay cuatro números primos —2, 3, 5 y 7— menores o iguales que 10], el teorema dice que a medida que aumenta N , el cociente $N/P(N)$ se va acercando cada vez más al logaritmo natural de N . Usando este resultado podemos calcular, por ejemplo, que aproximadamente el 3,6% del primer billón de números son primos. En mi opinión, éste es uno de los teoremas y demostraciones de la teoría de los números que exhiben una pureza inmutable tal que, al menos en los momentos en que uno está suficientemente meditabundo, parecen casi divinos.

Otro aspecto cuasidivino de los números primos es la facilidad con que se pueden formular enunciados relativos a ellos y cuya verdad o falsedad se desconoce. Un ejemplo es la conjetura de Goldbach, que dice que cualquier número par mayor que dos es la suma de dos primos. Comprobamos que $4 = 2 + 2$, $6 = 3 + 3$, $8 = 3 + 5$, $10 = 5 + 5$, $12 = 5 + 7$, y así sucesivamente, pasando por $374 = 151 + 223$, etc. Se cree que la conjetura es correcta, pero no se ha demostrado. Tampoco se sabe si hay o no una infinidad de pares de números primos que, como 17 y 19, 41 y 43 o 59 y 61, difieran en 2. Como la anterior, se cree que la proposición es verdadera pero todavía no ha sido demostrada.

A pesar de la metáfora de su divinidad, estas reflexiones puras y aparentemente inútiles acerca de los primos tienen su importancia para las tarjetas de crédito, las telecomunicaciones y la seguridad nacional. La idea básica consiste en que, aunque encontrar el producto de dos primos de 100 dígitos sea algo elemental, es prácticamente imposible factorizar el número de 200 dígitos que se obtiene. Esta propiedad de los números primos y otras parecidas, así como los anteriores números monstruosamente largos, se pueden utilizar para codificar mensajes que sólo podrá descifrar alguien que conozca de antemano sus factores primos. Los bancos usan a diario esos códigos para transferir fondos, y la National Security Agency utiliza variantes de los mismos para fines militares y de información. Tales aplicaciones pueden parecer tan absurdas como un mono trabajando en una fábrica de municiones. ¡Bendito Euclides!

P.D.: Hasta 1990 el mayor número primo conocido es el $[(391.581 \times 2^{216.091}) - 1]$, Para determinarlo hizo falta que un superordenador trabajara sin parar durante más de un año.

Números racionales y números irracionales

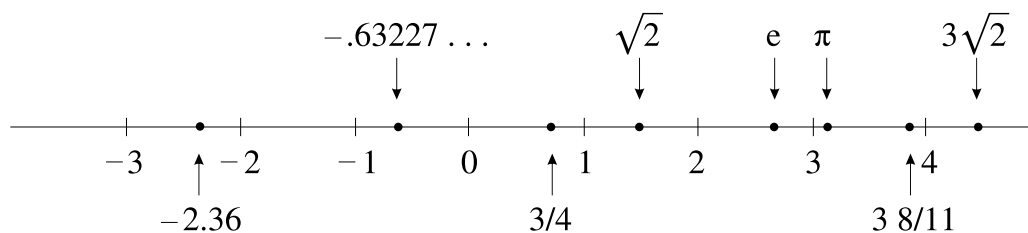
Si usted no es matemático ni tiene ninguna relación con la matemática, las definiciones de número racional y número irracional no le impresionarán demasiado, al menos de entrada. Número racional es aquél que se puede expresar como cociente de dos números enteros (como las fracciones), mientras que irracional es el número que no admite una expresión de este tipo. La mayoría de números que uno encuentra en la vida cotidiana son racionales: 3, que se puede expresar como $3/1$, 82, que se escribe $82/1$, $4\frac{1}{2}$ que se expresa $9/2$, $-17\frac{1}{4}$ que es $-69/4$, 35,28 que es $3.528/100$, o 0,0089 que es $89/10000$. Por otra parte, entre dos números racionales cualesquiera, por cerca que estén uno de otro, hay muchos más. Si la gente se pusiera a pensar en ello, probablemente llegaría a la conclusión de que el adjetivo «racional» en «número racional» es tan redundante como el adverbio «arriba» en «subir arriba». Sin embargo, los números racionales, como la vida misma, están flotando en un mar de irracionalidad, y en un sentido importante y bien definido debido al matemático Georg Cantor (véase la entrada sobre *Conjuntos infinitos*), hay muchos más números irracionales que racionales. Todos ellos, racionales e irracionales por igual, constituyen lo que se conoce como números reales y se pueden expresar en forma decimal y ordenar sobre una línea que se denomina, con bastante propiedad, la recta real.

—El primer número irracional que se descubrió fue la raíz cuadrada de 2, $\sqrt{2}$ (el número que multiplicado por sí mismo da exactamente igual a 2), y el descubrimiento de su irracionalidad originó una cierta crisis en la matemática de la antigua Grecia. Igual que ocurre con la demostración de la infinitud de los números primos, la demostración indirecta estándar de la irracionalidad

de $\sqrt{2}$ es tan elegante e ilustrativa de la técnica clásica de reducción al absurdo que la presentaré aquí. Supongamos que, contra lo que creemos, $\sqrt{2}$ fuera racional e igual al cociente de P entre Q o P/Q. Simplifiquemos los factores comunes de P y Q y reduzcamos la fracción a su mínima expresión. Por ejemplo, si la fracción fuera 6/4 podríamos reducirla a 3/2. Escribiremos esta fracción reducida como M/N, donde M y N son enteros que no tienen factores comunes.

Elevando al cuadrado ambos miembros de la ecuación $\sqrt{2} = M/N$, obtenemos $2 = M^2/N^2$ que, multiplicada por N^2 , se convierte en $2N^2 = M^2$. Veamos ahora cómo esto nos lleva inevitablemente a una conclusión absurda (y de ahí la reducción al absurdo) con lo que se demuestra la insostenibilidad de nuestra suposición previa de la racionalidad de $\sqrt{2}$. Esta sensación de inevitabilidad constituye, en mi opinión, una buena parte de la recompensa psicológica del trabajo en matemáticas y debería valorarse aunque no se apreciaran plenamente los detalles de la demostración.

Sigamos sin más dilaciones. Observamos que el miembro de la izquierda de la ecuación $2N^2 = M^2$ tiene un factor 2 y, por tanto, tiene que ser par. También habrá de serlo pues el miembro de la derecha. Como M^2 es par, también lo será M, pues el cuadrado de un número impar es impar. Así pues, al ser M par será igual a 2K para algún entero K, con lo que $M^2 = (2K)^2 = 4K^2$. Sustituyendo M^2 en la primera ecuación, tenemos $2N^2 = 4K^2$, y dividiendo por 2, $N^2 = 2K^2$. Como el miembro de la derecha de esta última ecuación tiene un factor 2, y por tanto es par, también lo han de ser N^2 y N, pues el cuadrado de un número impar es impar. Como N es par, se podrá escribir como 2J para algún entero J. Hemos insistido al principio en que M y N no tenían factores comunes, pero, como hemos demostrado, M y N tienen un factor 2 en común, por ser M igual a 2K y N a 2J. Esta contradicción es consecuencia directa de nuestra suposición original de que $\sqrt{2}$ es un número racional y, en consecuencia, dicha suposición es insostenible. En conclusión, $\sqrt{2}$ ha de ser irracional. Y esto es todo. QED. Sonido de trompetas.



Recta real en la que se señalan algunos números racionales en la parte inferior y algunos irracionales en la parte superior

Podemos demostrar la irracionalidad de muchos otros números.—El producto de dos números racionales es racional, y como sabemos que $\sqrt{2}$ es irracional, también habrá de serlo $\sqrt{2}/2$. De lo contrario $2 \times \sqrt{2}/2$ (que es igual a $\sqrt{2}$) sería también racional y acabamos de demostrar que no lo es. Por la misma razón son irracionales $\sqrt{2}/3$, $\sqrt{2}/4$, $\sqrt{2}/5$, etc. Otros números algebraicos, como la raíz cuadrada de 3, la raíz cúbica de 5 o la raíz séptima de 11 son también irracionales, al igual que π , e y una horda innumerable de números anónimos. (A efectos de cálculo nos basta con la aproximación racional de estos números irracionales. Para aproximar $\sqrt{2}$ podemos usar 1,4 o $14/10$, y, si necesitamos una mayor precisión, usaremos 1,41 o $1,414$).

Cuando escribimos $\sqrt{2}$ (cuya expresión decimal sigue 1,414 213 562 373 095 048 801 688 724 209 698 078 569 7...) o cualquier otro número irracional en forma decimal, encontramos que su desarrollo infinito no consiste en un grupo de cifras que se repite periódicamente. Por el contrario, los números racionales, todos ellos de hecho, tienen sucesiones de dígitos que se repiten. Los decimales 5,33333..., 13,87500000..., 29,38 461538 461538 461538..., que representan $5 \frac{1}{3}$, $13 \frac{7}{8}$, y $29 \frac{5}{13}$, respectivamente, son todos números decimales con cifras decimales que se repiten. Como cualquier número racional se puede expresar como cociente de dos enteros, la razón de esta repetición es clara. El primer residuo del proceso de división sólo puede ser un número menor que el divisor y, por tanto, sólo puede tomar un número finito de valores. Si repetimos el proceso indefinidamente, en algún momento habrá de volverse a repetir el mismo residuo, y a partir de este instante es como si el proceso volviera a empezar, de ahí la pauta que se repite en la expresión decimal de un número racional. La demostración del enunciado inverso es más delicada, pero se puede demostrar que si hay una

pauta repetitiva en el desarrollo decimal de un número, entonces éste se puede expresar como la suma de una serie geométrica infinita que siempre da un resultado racional.

Repito, un número es racional si y sólo si su desarrollo decimal se repite a partir de un cierto lugar. Las expresiones decimales de $\sqrt{2}$, π y e no presentan dicha repetición. En el conjunto de todas las expresiones decimales (es decir, en el conjunto de todos los números reales) es mucho más raro que haya una pauta y una repetición que la ausencia de las mismas. La armonía es siempre mucho más rara que la cacofonía.

Para terminar, he de señalar que, a pesar de su rareza, los números racionales tienen en muchos asuntos prácticos un papel mayor que los irracionales. En la mayoría de asuntos económicos cotidianos la distinción entre unos y otros es menos importante que la habilidad para manejar con soltura las operaciones con números racionales: fracciones —propias, impropias y mixtas—, decimales y porcentajes. Desgraciadamente, esta capacidad no es tan universal como convendría. La mayoría sabe manejar números racionales amables como 325,84 ptas. [o $32584/100$ ptas.], pero a muchos les supone un gran aprieto decidir si el número racional $25/3 \times [(8/9) - (2/5)]$ ptas. es mayor o menor que el número racional 4 ptas. [$4/1$ ptas.]. Y, por experiencia personal, sé que en Estados Unidos a pocos cajeros de supermercado les hace gracia que les llegue un cliente con un carro lleno de sandías pretendiendo alegre y enfervorecidamente que ha de bastar con un cuarto de dólar al precio anunciado de 0,59 centavos la libra [$(59/100)$ de centavo o $(59/10000)$ de dólar].

Ordenaciones parciales y comparaciones

La mayoría de fenómenos interesantes para el hombre se describen mejor mediante la estructura matemática conocida como ordenación parcial que con otros tipos de ordenación. Una ordenación parcial es una ordenación de un conjunto (esto es, algunos elementos del conjunto son mayores que otros) en la que algunos pares de elementos no son comparables. Es distinta de una ordenación lineal o total como «es tanto o más alto que» o «pesa tanto o más que». En estos dos casos está claro que dadas dos personas cualesquiera, una es más alta o pesa más que la otra. En una ordenación parcial puede suceder que los dos elementos simplemente no sean comparables con respecto a esa relación de orden.

Consideremos, a modo de ejemplo, el conjunto de círculos contenidos en un plano. Cada uno de ellos contiene y está contenido en otros círculos, pero si tomamos dos círculos al azar lo más probable es que ninguno de los dos contenga al otro. La mayoría de pares de círculos no son comparables, con lo que la relación «contiene» es una ordenación parcial. Cualidades tales como la belleza, la inteligencia y hasta la salud se pueden discutir con menos simplismo en términos de una ordenación parcial que según una ordenación lineal o total.

De hecho, la raíz de muchos problemas estriba en el intento de convertir un orden parcial en un orden total. Reducir la inteligencia a un orden lineal —un número en una escala de CI— ejerce una violencia sobre la complejidad y la diversidad de los talentos de las personas. Y lo mismo ocurre con los índices de belleza o de salud. Si intentamos hacer valer nuestras preferencias entre una gran opción de candidatos políticos topamos con varias paradojas de voto (véase la entrada sobre *Sistemas de votación*) y la expresión

«espectro político» es simplista y reduccionista.

Dificultades parecidas se plantean si clasificamos a nuestros amigos. La mayor parte de temas personales y públicos con que nos enfrentamos es lo suficientemente complicada y multidimensional como para que cualquier intento de hacerlos entrar en una lista de Procusto sea un indicio de miopía y de estrechez de miras (esto último también literalmente).

No obstante, las listas son tentadoramente simples. La gente siempre quiere saber quién es el «mejor» en tal o cual especialidad médica (la relación entre el sexismo y las jerarquías lineales no es casual), quien gana más dinero, quien va a la cabeza en la lista de títulos más vendidos. A veces me pregunto si quienes enfocan la vida más equilibrada y armoniosamente no estarán en desventaja en un mundo dominado por la obsesión y la monomanía. Quizás un modo de obsesión limitada o controlada sea la respuesta apropiada, aunque esto suena a oxímoron.

(Se habla mucho acerca de las diferencias en el rendimiento matemático de los hombres y las mujeres y, en particular, del número de personas de cada sexo que siguen estudios de matemática superior. No hay ninguna prueba convincente de que tales diferencias tengan una base genética. Mi sospecha es que se deben a la socialización y, *quizás*, a diferencias de personalidad influidas por la genética. Posiblemente la siguiente observación particular sea relevante al respecto. Aunque conozco a unas cuantas programadoras de primera clase, me he encontrado con muy pocas que se dediquen a la piratería^[8] y sean capaces de pasarse la noche en vela escribiendo programas que no sirven para nada, que lleven el pelo y la ropa sucios y tengan bolsas bajo unos ojos vidriosos de tanto mirar a la pantalla del ordenador, que no tengan amigos, que subsistan a base de patatas fritas, bocadillos y cafés, que cambien las configuraciones de sus sistemas a cada hora y que acostumbren a desaparecer detrás de imperios electrónicos creados por ellas mismas. Ya sé que puede sonar a condescendiente, pero las programadoras que conozco son generalmente demasiado equilibradas para ser piratas informáticos. Quizá la investigación matemática, sin llegar al grado de apasionamiento y de inmadurez monomaniaca de esa piratería, sea algo que atraiga con mayor facilidad a individuos con rasgos de personalidad que tradicionalmente se atribuyen al sexo masculino).

Volviendo a los órdenes parciales, estoy convencido de que, para comparar cosas, a menudo nos sirve mejor un árbol o un arbusto que una barra. Los árboles y los arbustos permiten encajar tanto los elementos incomparables (sobre ramas distintas) como los comparables (sobre una misma rama), mientras que las barras lo reducen todo a una sola dimensión.

Oulipo: matemática en la literatura

El Ouvroir de Literature Potentielle (Taller de literatura potencial), abreviado Oulipo, es el nombre de un reducido grupo de escritores, matemáticos y académicos, originariamente franceses, que se dedica a la exploración de técnicas matemáticas y cuasimatemáticas aplicadas a la literatura. Fue fundado en París en 1960 por Raymond Queneau y François Le Lionnais, y busca nuevas estructuras literarias por imposición de condicionantes poco corrientes, métodos sistemáticos de transformar textos y maneras de ejemplificar conceptos matemáticos mediante palabras. El grupo ha publicado varios manifiestos pero, igual que sucede con los entusiastas del yoga, son más interesantes sus técnicas y sus resultados que su filosofía.

Los *Cien billones de sonetos* de Queneau es un ejemplo fundamental de la aproximación combinatoria de Oulipo a la literatura. La obra consiste en diez sonetos escritos en diez páginas, las cuales están cortadas de modo que cada uno de los catorce versos de cada soneto se pueda tomar independientemente de los demás. Así, cada uno de los diez primeros versos se puede combinar con cualquiera de los diez segundos versos, con lo que se tienen 10^2 o 100 pares distintos de versos iniciales. Cualquiera de estas 100 posibilidades se puede combinar con cualquiera de los diez terceros versos, obteniendo 10^3 o 1.000 ternas de versos posibles. Reiterando el procedimiento nos encontramos con que hay 10^{14} sonetos posibles. (Véase la entrada sobre *La regla del producto*). Queneau sostiene que todos ellos tienen sentido, aunque seguramente nadie lo comprobará, pues estos 10^{14} sonetos representan muchísimos más textos que todo el resto de la literatura mundial.

Otro buen ejemplo del trabajo de Oulipo es el algoritmo $(N + 7)$ de Jean

Lescure para transformar un texto. Se toma un extracto de un periódico, una novela o un libro sagrado concreto y se sustituye cada nombre que aparezca en él por el séptimo nombre que le sigue en un diccionario corriente. Si el original está bien escrito, el texto resultante normalmente conserva el ritmo y algunas veces hasta parte del sentido original. «En el priorato díptero creó los cienos y la tiesura. Y la tiesura estaba desordenada y vacía, y los tinteros estaban sobre la ablución...». El algoritmo se puede modificar, naturalmente, tomando el décimo nombre que siga en el diccionario a la palabra sustituida, etc.

Y otra obra más, esencial también en el trabajo de Oulipo, la novela de 300 páginas de Georges Perec, *La Disparition*, no contiene una sola letra «e» aparte, claro está, de las cuatro desafortunadas veces que aparece en el nombre del autor. Es para pensarlo, sin ningún «el», ni «donde», ni «ella», ni «porque», ni tan siquiera un «pero». En un ensayo sobre tales lipogramas, obras que prescinden de determinadas letras, Perec defiende la sensatez y la seriedad de tales empresas, y aduce que los condicionantes y la artificiosidad han sido los motores que han llevado, no sólo a los miembros de Oulipo, sino también a muchos grandes autores (François Rabelais, Laurence Sterne, Lewis Carroll, James Joyce, Jorge Luis Borges y el miembro de Oulipo, Italo Calvino, entre otros) a sondear todas las posibilidades de un idioma.

Algunas de esas posibilidades tienen que ver con los modos de unir un texto con otro «multiplicándolos» (como las matrices y los números), con encontrar la intersección lógica de dos obras dispares, o con «haikuificar» un poema largo convirtiéndolo en otro más corto. Cuentan también con las formas más corrientes de juegos de palabras. Los palíndromos, frases escritas que se leen igual al derecho que al revés (el ejemplo clásico en español: «Dábale arroz a la zorra el abad»); las transposición de sonidos en dos o más palabras (llamado en inglés *spoonerisms*: «Chuck you, Farlie»);^[9] los preverbios del escritor americano miembro de Oulipo, Harry Mathews, que funde dos proverbios en uno («A rolling stone gets the worm», o «A bird in the hand waits for no man»);^[10] las frases tipo bola de nieve, en las que cada palabra tiene una letra más que la anterior («I do not pass Sally unless feeling reckless»)^[11] y todos sus parientes próximos y lejanos son omnipresentes en

los poemas, cuentos y novelas de Oulipo.

Aunque pueda parecer raro, Oulipo no ha manifestado hasta muy recientemente el interés por los ordenadores que podría desprenderse de su inclinación a la permutación sintáctica. Adaptando varios programas de tratamiento de textos con diccionarios especiales, de sinónimos, contadores de palabras clave y lo que se ha dado en llamar utillajes de hipermedia, se podría facilitar su juego literario combinatorio. (Véase la entrada sobre *La conciencia humana y su naturaleza fractal*).

Las obras de Oulipo, cuando son buenas, son refrescantes, estimulantes y divertidas. Si no, su excesiva truculencia y sus incesantes juegos de palabras son tan pesados como cuidar a un niño locuaz y listillo. (Obsérvese que he logrado la asombrosa proeza de escribir toda la entrada con sólo veintiséis letras: no he usado para nada la «ñ»).

La paradoja de Russell

En la literatura moderna y en el cine es cada vez más corriente que aparezcan personajes que se salen del relato para comentarlo y, a veces, incluso para comentar sus comentarios. Que esta estratagema se utilice últimamente con más frecuencia puede deberse a un aumento de la timidez, a un sentido de la realidad más fragmentado y menos unitario, o a una mayor apreciación por las obras abstractas (véase la entrada sobre *Humor*). Sea cual sea la razón, se trata de una idea muy antigua. Así lo demuestran el coro del teatro griego clásico o sus distintas encarnaciones en el teatro medieval y en Shakespeare, que actuaban como comentaristas institucionalizados y al mismo tiempo jugaban un papel esencial en la obra.

Dicho comentario en varios niveles comporta una complejidad más viva y lleva a veces, como la clásica paradoja del mentiroso ya conocida hace 2.000 años, a situaciones paradójicas. Se dice que Epiménides el Cretense dijo que todos los cretenses son unos mentirosos. Lo esencial de esta declaración tan hipócrita queda más claro si simplificamos la afirmación y la convertimos en «Estoy mintiendo» o, mejor aún, «Esta frase es falsa». Si llamamos Q a «Esta frase es falsa», observamos que si Q es verdadera, entonces, por lo que ella misma dice, ha de ser falsa. Y por otra parte, si Q es falsa, entonces dice la verdad y, por tanto, Q ha de ser verdadera. Así pues, Q es verdadera si, y sólo si, es falsa. Una variante atenuada de la paradoja del mentiroso se da, aunque implícitamente, siempre que el marco de un cuadro, el escenario de una obra de teatro o el tono de voz que se emplea para contar un chiste sugieren, «Esto es falso, no es real».

La faceta de autorreferencia de estas paradojas se puede expresar también de otros modos. Consideremos el conocido caso del barbero de una ciudad a

quien la ley manda afeitar a todos los hombres pero no a aquellos que se afeitan a sí mismos. El pobre barbero se queda con la duda de si debe afeitarse o no. Si se afeita a sí mismo contraviene la ley. Pero por otra parte, si no lo hace también la desobedece. Otra versión de la paradoja, más próxima al tema central de esta entrada, trata de las leyes de residencia para los alcaldes de un cierto país. Algunos de los alcaldes viven en las ciudades que gobiernan y otros no. Un buen día, un monarca reformista publica un decreto que obliga a todos los alcaldes no residentes, y sólo a ellos, a vivir en un mismo lugar —lo llamaremos ZAD, de Zona de Alcaldes Desplazados—. De pronto, al darse cuenta de que ZAD ha de tener alcalde, el rey coge un real dolor de cabeza de tanto darle vueltas a dónde habría de residir el alcalde de ZAD.

Llegamos por fin a la paradoja de Russell. Se debe al filósofo y matemático inglés Bertrand Russell y, en cierto modo, tiene el sabor de las paradojas anteriores, aunque con una potencia matemática muy superior. Aunque está formulada en los términos de la teoría de conjuntos, lo único que hay que saber de los conjuntos es que son, al menos informalmente, colecciones bien definidas de objetos de cualquier clase. También es útil conocer un poquito de notación. Para indicar que 7 pertenece al conjunto P, el conjunto de los números primos, o que 10 no pertenece al conjunto P, se escribe $7 \in P$ y $10 \notin P$, respectivamente. Así, $13 \in P$ y $15 \notin P$. En general « $X \in Y$ » significa que X pertenece al conjunto Y o, lo que es lo mismo, que Y contiene a X como elemento. Por tanto, si Y es el conjunto de países de las Naciones Unidas y X es Kenia, entonces $X \in Y$.

Ahora bien, para llegar a la paradoja observaremos que algunos conjuntos se contienen a sí mismos como elementos ($X \in X$) y otros no ($X \notin X$). El conjunto de todas las cosas mencionadas en esta página es a su vez mencionado en esta página y, por tanto, se contiene a sí mismo como elemento. Análogamente, el conjunto de todos los conjuntos que tienen más de once elementos contiene a su vez más de once elementos y pertenece por tanto a sí mismo. Y el conjunto de todos los conjuntos es también un conjunto y por tanto también pertenece a sí mismo. Sin embargo, la mayoría de conjuntos que se dan de un modo natural no se contienen a sí mismos como elementos. El conjunto de senadores de la Cámara Alta no es un

senador, con lo que no pertenece a sí mismo. O el conjunto de los números impares no es un número impar y por tanto no se contiene a sí mismo como elemento.

Denotemos por M el conjunto de todos los conjuntos que pertenecen a sí mismos y por N el conjunto de todos los conjuntos que no tienen esta propiedad. Puesto en una forma más simbólica, para cualquier conjunto X , tenemos que $X \in M$ si y sólo si $X \in X$. Por otra parte, para cualquier X , $X \in N$ si y sólo si $X \notin X$. Ahora bien, podríamos preguntarnos si N pertenece a sí mismo o no. (Compárese esta pregunta con «¿Quién afeita al barbero?», o con «¿Dónde vive el alcalde de ZAD?»). Si $N \in N$, entonces por definición $N \notin N$. Pero si $N \notin N$, entonces por definición $N \in N$. Así pues, N pertenece a sí mismo si y sólo si no pertenece a sí mismo. Esta contradicción se conoce como paradoja de Russell.

La solución de Russell a esta paradoja (hay otras) consiste en restringir el concepto de conjunto a una colección bien definida de conjuntos ya existentes. En su famosa teoría de los tipos, él y Alfred North Whitehead clasificaron los conjuntos según sus tipos. En el nivel más bajo, tipo 1, están los objetos individuales. En el siguiente nivel, tipo 2, están los conjuntos de objetos de tipo 1. En el siguiente nivel, tipo 3, están los conjuntos de objetos de tipo 1 o de tipo 2, etc. Los elementos de tipo N son conjuntos de objetos de tipo $(N - 1)$ o inferior. Así se evita la paradoja porque un conjunto sólo puede pertenecer a un conjunto de un tipo superior y por tanto no puede pertenecer a sí mismo. Se ha eliminado la posibilidad de que un conjunto pertenezca a sí mismo ($X \in X$), y así carecen de sentido los conjuntos M y N definidos en base a este concepto.

(Tengo que decir que el objetivo de Russell y Whitehead al construir la teoría de los tipos era, además de evitar esta paradoja y otras parecidas, dar una base axiomática a toda la matemática. Consiguieron reducir toda la matemática a la lógica encarnada en la teoría de los tipos —la lógica más la anterior definición jerárquica de conjunto—. Quiero insistir también en que la referencia a uno mismo no implica necesariamente una paradoja de la matemática o de los lenguajes corrientes y, de hecho, normalmente no es así. La mayoría de casos en que se emplea la palabra «yo», por ejemplo, no dan en absoluto ningún problema).

Podemos también resolver la paradoja del mentiroso en términos de tipos y asignar a «Todos los cretenses son unos mentirosos» un tipo superior al de las otras frases que puedan pronunciar los cretenses. Distinguiremos entre frases de primer orden, que no hacen referencia a ninguna otra frase (verbigracia, «Llueve» o «Menelao es calvo»); frases de segundo orden, que pueden referirse a frases de primer orden («Los comentarios de Waldo acerca de la comida eran disparatados»); frases de tercer orden, que pueden referirse a las de segundo orden; y frases de todos los órdenes superiores. Así, cuando Epiménides dice «Todos los cretenses son unos mentirosos», se ha de interpretar que pronuncia una frase de segundo orden, que no se aplica a sí misma sino sólo a las frases de primer orden. Si dice que todas sus frases de segundo orden son falsas, esta afirmación es de tercer orden y no se aplica a sí misma.

Con más generalidad, esta estructura jerárquica se puede exportar al concepto de verdad: verdad₁ para los enunciados de primer orden, verdad₂ para los enunciados de segundo orden, etc. El matemático Alfred Tarski ha desarrollado y formalizado a fondo este concepto de verdad y el filósofo Saul Kripke la ha flexibilizado para poder aplicarla a los lenguajes naturales, más propensos a la confusión. Una de estas confusiones se produce cuando dos o más personas hacen afirmaciones inocuas que, al considerarlas juntas, dan lugar a una paradoja. Un ejemplo sencillo es la siguiente conversación en la que Jorge dice «Marta siempre miente» y Marta dice «Jorge siempre dice la verdad», y nadie dice nada más.

Lo que me resulta extraño es que, normalmente, la misma gente que considera que estos temas son imposiblemente esotéricos y académicos sea la que es capaz de desenmarañar rutinariamente narraciones muy barrocas y con una enrevesada superposición de capas, historias de intrigas adolescentes (ella dijo que él dijo, pero no puede ser cierto, pues ella no habría dicho esto y aquello cuando él le mintió acerca de lo que él pensaba que ella decía), de hacer análisis laberínticos de la política de Oriente Medio, o interesarse por las películas y novelas de las que hablé al principio. (Véase también la entrada sobre *Variables*).

Pi

Una vez hice una investigación informal entre mis amigos y vecinos no matemáticos para ver cuántos de ellos sabían qué era pi. Casi todos sabían que tenía algo que ver con la geometría, y en concreto con los círculos, siendo esto último una explicación indudable de por qué algunos de ellos creían que se escribía «pie».^[12] (La letra griega π , en español pi, se usa para denotar el conocido número desde el siglo XVIII). Unos pocos sabían que valía aproximadamente $22/7$ (tengo que decir en su honor que es una aproximación bastante buena), pero la mayoría daban estimaciones bastante alejadas (un eminente abogado llegó a decir con gran sonoridad que era 5,42). Algunos decían que π era el área del círculo, mientras que la mayoría trataron de disimular su ignorancia y/o mi impertinencia con una broma (el abogado me pidió que le recitara las estipulaciones de una ley de bancarrota). Sólo unos pocos sabían que era el cociente de la longitud de la circunferencia entre el diámetro. Esto es, π es igual a C/D , la circunferencia dividida por el diámetro.

Entre las estimaciones antiguas de π tenemos 3 (Antiguo Testamento: un problema insuperable, según parece, para los autores bíblicos), $25/8$ (babilonios), $256/81$ (egipcios), $22/7$ (griegos), $355/113$ (chinos, con seis cifras decimales correctas) y $\sqrt{10}$ (indios, una grata coincidencia). Una estimación mucho más precisa es 3,141 592 653 589 793 238 462 643 383 279 502 884 1972, pero π es un número irracional (no se puede expresar como cociente de dos números enteros), lo que implica que su expresión decimal tiene una longitud infinita y sin repeticiones periódicas. Es un número trascendente, y esto significa que no es la solución de ninguna

ecuación algebraica. Como una de las principales constantes de la matemática, π figura en muchas fórmulas importantes, y una fundamental es la del área del círculo, que es π veces el cuadrado de su radio ($A = \pi R^2$). Por contra, el volumen de la esfera vale $4/3$ por π por el cubo del radio ($V = 4/3 \times \pi \times R^3$). Pi es indispensable para la formulación de las célebres leyes del electromagnetismo del físico escocés James Clerk Maxwell y aparece también en otras muchas fórmulas y contextos donde su presencia es más sorprendente, pues no tienen nada que ver con círculos ni esferas que expliquen su aparición. Por ejemplo, $\pi/4 = 1 - 1/3 + 1/5 - 1/7 + 1/9 - 1/11 + \dots$ y $\pi^2/6 = 1/1^2 + 1/2^2 + 1/3^2 + 1/4^2 + 1/5^2 + 1/6^2 + \dots$

El problema de la aguja de Buffon, planteado por primera vez en el siglo XVIII por el conde de Buffon, tampoco tiene que ver con círculos y, sin embargo, su solución depende de π , y de hecho se usó una vez para calcular su valor. Supongamos con Buffon que tenemos un suelo hecho de tablas de madera de 6 centímetros de anchura. Supongamos también que tenemos una aguja de 6 centímetros y que la dejamos caer descuidadamente al suelo. ¿Cuál es la probabilidad de que la aguja caiga de modo que cruce una de las líneas que separan dos tablas adyacentes?

O lo que es lo mismo, ¿cuál es la probabilidad de que la aguja al caer no esté contenida en una sola tabla de madera? Me saltaré la deducción, pero se puede demostrar que la probabilidad es $2/\pi$.

Este resultado se podría usar para hacer una estimación de π del modo siguiente (véase además *El método de simulación de Montecarlo*): déjese caer la aguja al suelo 10.000 veces (ya sea porque uno se dedica a la matemática experimental o porque está en la cárcel y no tiene nada mejor que hacer) y determínese que fracción de las veces cae sobre una línea. Supongamos que esto ocurre en 6.366 casos. Nuestra estimación de la probabilidad de que la aguja caiga sobre una línea será pues de 0,6366. (La aguja caerá sobre una línea con esta frecuencia en el supuesto de que todos los puntos del suelo son igualmente probables y que todas las orientaciones de la aguja son también equiprobables). Si igualamos esta probabilidad a $2/\pi$ y despejamos π de la ecuación resultante, llegamos a 3,1417 como estimación de π , que es un resultado bastante próximo al valor real. No hace falta decir que hay métodos

incomparablemente mejores para determinar el valor de π .

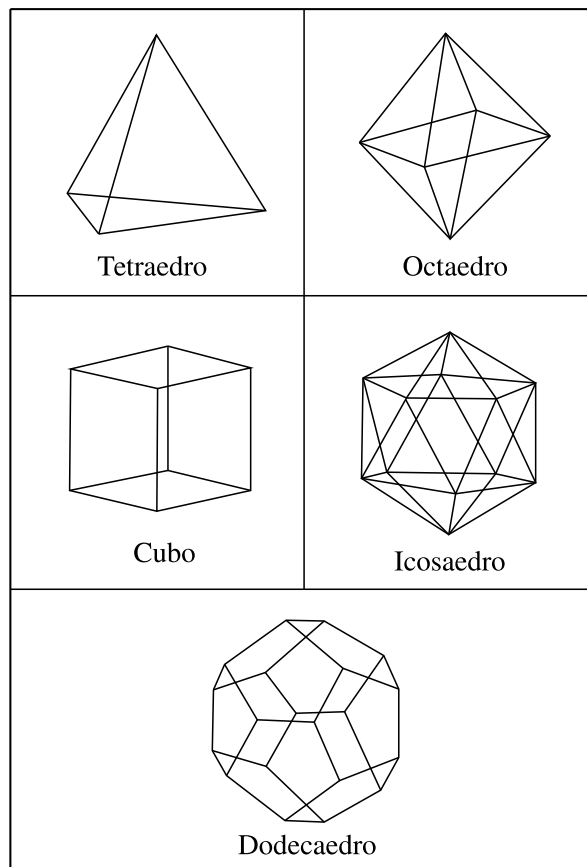
El atractivo de π reside, según mi opinión, en su universalidad. Como el Everest, es un desafío que siempre está ahí, frente a nosotros. Los últimos cálculos de π por superordenador (hacia 1990) dan más de mil millones de dígitos. Expresado en un sistema de numeración menos antropocéntrico, como el binario (en base 2), π podría incluso servir, como han sugerido muchos relatos de ciencia ficción, para comunicar nuestra sofisticación tecnológica y nuestra acogedora naturaleza euclídea a seres de otra galaxia. Aun a riesgo de que los más anuméricos de ellos pudieran interpretar la señal como una especie de partitura interestelar, el mensaje sería recibido.

Y para terminar, un pequeño problema: imaginemos una cuerda que da la vuelta alrededor de la Tierra a ras de suelo sobre el ecuador. ¿Cuánta cuerda habrá que añadir para que la cuerda alargada dé la vuelta a la Tierra un metro por encima del suelo siguiendo también el ecuador?

La respuesta, que al principio sorprende, es que basta con algo más de seis metros de cuerda. La explicación se basa en la fórmula del perímetro de la circunferencia: $C = \pi D$. Si el diámetro de la Tierra es aproximadamente 13.000 kilómetros, la longitud de la cuerda es $\pi \times 13.000$ kilómetros, que da aproximadamente 40.820 kilómetros. Si estipulamos que la cuerda vaya ahora un metro por encima del suelo siguiendo el ecuador, estamos pidiendo que el diámetro aumente en dos metros. Por tanto, la segunda circunferencia medirá $\pi \times (13.000 \text{ kilómetros} + 2 \text{ metros})$, lo cual da $(\pi \times 13.000 \text{ kilómetros}) + (\pi \times 2 \text{ metros})$, es decir, aproximadamente igual a 40.820 kilómetros más 6,28 metros. Así pues, sólo habremos de añadir 6,28 metros de cuerda.

Los poliedros regulares

En cierta ocasión tuve una estudiante que para referirse a los sólidos platónicos decía las ensaladas de César^[13] Por lo demás esa estudiante era más bien mediocre, y nunca llegué a saber si se trataba de un chiste intencionado o una manifestación de «angustia matemática». En cualquier caso los poliedros regulares o sólidos platónicos son cuerpos tridimensionales cuyas caras son polígonos regulares congruentes que en los vértices forman ángulos iguales. Un cubo es un poliedro regular porque todas sus caras son cuadrados iguales, mientras que una caja de zapatos no lo es porque no todos los lados son congruentes. Otro ejemplo de sólido regular (o platónico) es el tetraedro, cuyas cuatro caras son triángulos equiláteros del mismo tamaño. (En algunos países los envases de leche de cartón son tetraédricos).^[14] No puedo poner demasiados ejemplos porque, como ya descubrieron los antiguos geómetras griegos, sólo hay cinco poliedros regulares.



Los cinco poliedros regulares

Los cinco son los ya citados tetraedro y cubo, el octaedro cuyas ocho caras son triángulos equiláteros iguales, el dodecaedro cuyas doce caras son pentágonos regulares (equiláteros y equiángulos) iguales y el icosaedro cuyas veinte caras son triángulos equiláteros iguales. Su belleza y corto número (¿por qué sólo cinco?) propició una reverencia mística hacia estos cinco sólidos regulares, una de cuyas más claras manifestaciones fue el intento de Johannes Kepler de explicar el movimiento de los planetas del sistema solar en base a las propiedades de dichos sólidos. (Afortunadamente no abandonó la tarea y, más tarde, descubrió unas leyes planetarias basadas en la observación y no en prejuicios matemáticos). Aún hoy en día algunas personas meditan en el interior de grandes estructuras tetraédricas o acarician cristales simétricos en la creencia de que por medio de ello alcanzarán conocimiento, salud o algún otro desiderátum. Es curioso que algunas de las

discusiones en torno a las formas de las siete discontinuidades elementales de la moderna teoría de catástrofes tengan este mismo carácter sobreexcitado y pseudocientífico.

Misticismos aparte, los poliedros regulares tienen algo de prístino. Que solo haya cinco se puede demostrar fácilmente a partir de una fórmula bien conocida debida al matemático suizo del siglo XVIII Leonhard Euler. Según esta fórmula, cualquier poliedro (un sólido cuyas caras son polígonos que no han de ser necesariamente iguales, ni regulares, ni tener la misma forma) cumple que, si se suma el número de vértices con el número de caras y al resultado se le resta el número de aristas, siempre da 2. Puesto en un lenguaje más matemático, y usando siglas de significado evidente, el resultado es $V + C - A = 2$. (Algo que ayuda a entender la fórmula es comprobar su validez para cubos y cajas de zapatos, en cuyo caso $V = 8$, $A = 12$ y $C = 6$). Como la fórmula es válida para todos los poliedros (ya sean platónicos o no) y como el hecho de ser platónico impone nuevas condiciones sobre V , E y F , un poco de cálculo demuestra que sólo son posibles los cinco poliedros regulares citados.

Sorprendentemente, en las fórmulas del área y del volumen de los poliedros regulares interviene el número e , la base de los logaritmos naturales. Hay también una relación inesperada entre el rectángulo áureo (véanse las entradas sobre E y *Rectángulo áureo*) y el icosaedro: Si tres rectángulos áureos iguales se cortan simétrica y perpendicularmente, sus vértices son los vértices de un icosaedro regular. Una propiedad más comprensible de los poliedros regulares es que nos proporcionan buenos dispositivos para el azar: los cubos (esto es, los dados) para elegir cualquiera de entre seis resultados equiprobables, los dodecaedros para cualquiera de entre doce, etc. De hecho, la naturaleza elemental de los poliedros regulares (al igual que la de π y e) les convierte en buenos candidatos para comunicarnos con otros seres inteligentes si los hubiera en alguna parte.

Los poliedros regulares son uno de los pocos temas interesantes de la geometría clásica de los sólidos, una materia que, moribunda en su tiempo, ha resucitado a través de sus generalizaciones a dimensiones superiores, y que actualmente está en la vanguardia de la investigación en algunas ramas de la física, del álgebra abstracta y de la topología.

Probabilidad

Todo el mundo tiene una idea intuitiva de la probabilidad. A veces tan primitiva como la del barbero que una vez me contaba su estrategia para la lotería: «Tal como yo lo veo, puedo ganar o perder, mis posibilidades son pues mitad y mitad». Sin embargo, a pesar de que hay muchos aspectos mal comprendidos (véase la entrada sobre *Coincidencias*), las afirmaciones viscerales de la gente relativas a la probabilidad suelen ser considerablemente más sofisticadas. Soltamos fácilmente frases como «la probabilidad de que salga cara», «la probabilidad de que Marta se case con Jorge» y «la probabilidad de que llueva mañana durante el partido», y la mayoría parece tener claro su significado. Sólo si preguntamos qué es efectivamente la probabilidad nos encontramos totalmente desconcertados y perplejos.

La pregunta no es fácil de contestar, aunque se haya dado una serie de respuestas tentativas. Algunos han concebido la probabilidad como una relación lógica, como si sólo con echar una mirada a un dado, apreciar su simetría y emplear métodos lógicos, se pudiera decidir que la probabilidad de que salga 5 ha de ser $1/6$. Otros han sugerido que la probabilidad es simplemente una cuestión de creencia subjetiva, nada más que una expresión de una opinión personal. Según otros, la clave del análisis es la frecuencia relativa, y la probabilidad de un suceso sería una manera abreviada de indicar el porcentaje de veces que se produce a largo plazo, aunque no suelen explicar qué significa «a largo plazo».

Hay aún otras variantes y otras versiones, pero ninguna de ellas es universalmente convincente. La historia ha acabado finalmente así: los matemáticos se han retirado y se han declarado victoriosos al mismo tiempo. Han observado que, como cualquier definición razonable de probabilidad ha

de tener ciertas propiedades formales, la probabilidad se define como aquello que tenga precisamente dichas propiedades. No es muy gratificante filosóficamente hablando, pero al menos es matemáticamente liberador.

Nótese que esto es parecido a lo que ocurrió en geometría con las rectas y los puntos. Euclides dio unas definiciones vacías de estos conceptos que en realidad nunca usó, mientras que otras aproximaciones más modernas a la geometría del plano, siendo axiomáticas, definen los puntos y las rectas como cualquier cosa que cumpla las propiedades indicadas por los axiomas. El matemático ruso A. N. Kolmogorov es el padre de esta formulación abstracta de la teoría de la probabilidad. En vez de describir su elegantemente escueto formalismo, presentaré algunas propiedades y teoremas fundamentales, la mayoría de los cuales son conocidos desde que, en sus orígenes en el siglo XVII, la teoría de la probabilidad iba de la mano de los juegos de azar.

Para empezar, la probabilidad es un número comprendido entre 0 y 1. El 0 indica imposibilidad, el 1, certeza, y los valores intermedios, grados intermedios de probabilidad. Equivalentemente, podemos tomar el dominio de valores entre el 0% y el 100%. Si dos o más sucesos son mutuamente excluyentes, la probabilidad de que ocurra alguno de ellos se obtiene sumando sus probabilidades individuales. Así, la probabilidad de que un terrestre elegido al azar sea chino, indio o norteamericano es aproximadamente del 45% (25% de que sea chino más el 15% de que sea indio más el 5% de que sea norteamericano).

Dados dos sucesos arbitrarios (con tres o más sucesos valen fórmulas análogas), la probabilidad de que ocurra al menos uno de ellos es algo más difícil de obtener: primero se suman las probabilidades individuales y luego se resta del resultado la probabilidad de que ocurran ambos a la vez. Si en un gran edificio de apartamentos de Nueva York el 62% de los inquilinos lee *The New York Review of Books*, el 24% lee el *National Inquirer* y el 7% lee ambas revistas, la probabilidad de que un inquilino tomado al azar lea al menos una de dichas revistas es del 79% ($62\% + 24\% - 7\%$). La probabilidad de que un suceso no se produzca es el 100% menos la probabilidad de que sí se produzca. Por tanto, una probabilidad del 79% de leer una de las revistas al menos significa una probabilidad del 21% de no leer ninguna de ellas.

El concepto de independencia tiene una importancia crucial en la teoría

de la probabilidad. Se dice que dos sucesos son independientes si el hecho de que se produzca uno de ellos no influye sobre la probabilidad de que se produzca el otro. Si lanzamos dos veces una moneda al aire, cada tirada es independiente de la otra. Si tiramos un par de dados, lo que sale en uno es independiente de lo que sale en el otro. Si escogemos dos personas del listín telefónico, la altura de una es independiente de la de la otra.

Calcular la probabilidad de que se produzcan dos sucesos independientes es cosa fácil: basta simplemente con multiplicar sus probabilidades respectivas. Así, la probabilidad de que salgan dos caras es $1/4$ ($1/2 \times 1/2$). La probabilidad de que al tirar dos dados salga 2, esto es, (1,1), es $1/36$ ($1/6 \times 1/6$), mientras que la de que salga 7 es $6/36$, pues hay seis maneras mutuamente excluyentes de que los números que salgan en los dos dados sumen 7 [(1,6), (2,5), (3,4), (4,3), (5,2), (6,1)] y cada una de ellas tiene una probabilidad de $1/36$ ($1/6 \times 1/6$). La probabilidad de que dos personas tomadas al azar de un listín telefónico midan más de 2 metros se obtiene elevando al cuadrado la probabilidad de que una sola persona escogida por el mismo procedimiento mida más de 2 metros.

Esta regla del producto de la probabilidad puede generalizarse y aplicarse en sucesiones de sucesos. La probabilidad de sacar un 3 con un dado cuatro veces consecutivas es de $(1/6)^4$; la de que, al tirar una moneda, salga cara seis veces seguidas es de $(1/2)^6$; la de que alguien sobreviva a tres disparos en la ruleta rusa es de $(5/6)^3$. Si tomamos un libro que sea juzgado positivamente sólo por el 10% de sus lectores (y el 90% lo encuentre abominable) y lo sometemos a una docena de críticos, la probabilidad de que todos y cada uno de ellos hagan críticas negativas es $(0,9)^{12}$, o 0,28, y por tanto, la probabilidad de que el libro guste al menos a uno de los doce es $1 - 0,28$, o 0,72. Así, incluso con un «mal» libro, la posibilidad de recoger unas cuantas críticas favorables aumenta con el número de críticos, lo cual, unido al esmero en extraer lo más conveniente de críticas poco entusiastas, nos da una explicación de los «estilo ágil y directo», «increíble fuerza», ... en las sobrecubiertas de los libros.

Naturalmente, ocurre a menudo que los sucesos no son independientes; el hecho de que se produzca uno hace que el otro sea más o menos probable. Si

hemos sacado un 6 con el primer dado, la probabilidad de que la suma de las caras de los dos dados sea 10, 11 o 12 es mayor que si no conociéramos dicho resultado. Si sabemos que una persona mide más de 2 metros, disminuye la probabilidad de que pese menos de 60 kilos. Si en un cierto vecindario hay un gran número de Mercedes, probablemente habrá pocas personas sin hogar en él. Estos pares de sucesos son todos dependientes.

Lo que nos interesa determinar en tales casos es la probabilidad condicional de que ocurra o haya ocurrido uno de los sucesos sabiendo que el otro se producirá o se ha producido ya. La probabilidad condicional de que la suma de los dados sea 10, 11 o 12 habiendo salido 6 en el primer dado es $1/2$. Hay seis posibilidades igualmente probables [(6,1), (6,2), (6,3), (6,4), (6,5), (6,6)] y tres de ellas suman 10 o más. Y me atrevería a decir que la probabilidad condicional de que uno pese menos de 60 kilos sabiendo que mide más de 2 metros no excede el 5%, considerablemente menos que la probabilidad de que una persona tomada al azar pese menos de 60 kilos.

Al tratar con probabilidades condicionales hay que ser muy cuidadoso. Nótese, por ejemplo, que la probabilidad condicional de que uno hable español sabiendo que tiene nacionalidad española es aproximadamente del 95%, mientras que la probabilidad condicional de que uno sea ciudadano español sabiendo que habla español no es mucho más del 10%. O considérese la escena siguiente, que es una aclaración de otra sacada de mi libro *El hombre anumérico* sobre la que he recibido un gran número de cartas. Se sabe que en cierto vecindario curiosamente «normal» de los años cincuenta vive una familia de cuatro personas en cada casa: el padre, la madre y dos hijos. Uno escoge una casa al azar, toca el timbre y abre una chica. (Supondremos que en los años cincuenta, la chica de la casa, *si* la hay, es la que siempre abre la puerta). Suponiendo lo dicho, ¿cuál es la probabilidad condicional de que esta familia tenga un hijo y una hija? La respuesta, que quizá pueda sorprender, no es $1/2$ sino $2/3$. Hay tres posibilidades igualmente probables —el mayor es chico y la menor, chica; la mayor es chica y el menor, chico; y las dos son chicas— y en dos de ellas hay un hijo en la familia. La cuarta posibilidad —dos chicos— queda descartada por la sencilla razón de que nos ha abierto una chica.

El recuento de probabilidades de sucesos complejos no es, en general,

difícil *si* nos dan las probabilidades de los sucesos simples que los constituyen. Podemos usar los axiomas de Kolmogorov (probabilidad de sucesos mutuamente excluyentes, sucesos independientes, etc.), descomponer los sucesos complejos en subsucesos mutuamente excluyentes y calcular. O, si esto resulta demasiado complicado, podemos simular la situación con un ordenador y determinar empíricamente la respuesta (véase la entrada sobre *El método de simulación de Montecarlo*).

La asignación de probabilidades a los sucesos elementales es, sin embargo, una tarea considerablemente más ardua. Existe el problema de que las percepciones de la gente acerca de la delincuencia o de la enfermedad, por ejemplo, se han formado más a partir de las escenas dramáticas de los telediaros que de las mismas estadísticas sobre delincuencia o salud. Sumemos las probabilidades de que cualquiera de los cinco mil millones y pico de habitantes del mundo le mate a usted. El resultado, tristemente alto en Estados Unidos, todavía es menor que la probabilidad de que usted se suicide. O bien considere que, dado un habitante medio de Estados Unidos, es un cuarto de millón de veces más probable que muera de una enfermedad cardíaca que de botulismo, intoxicación mortal por ingerir conservas en mal estado. No hace falta decir que un asesinato o un caso de botulismo son fácilmente noticiables, mientras que el suicidio y los ataques de corazón no lo son (a menos que se trate, naturalmente, de un personaje famoso). El problema no es meramente académico. La incapacidad de tasar los riesgos que nos acechan y de ponerlos en una perspectiva global lleva generalmente a una ansiedad personal paralizante e infundada o a demandas inasequibles y económicamente prohibitivas de un entorno libre de riesgos.

Y, sin embargo, incluso cuando calculamos y estimamos probabilidades en la más ideal de las situaciones, permanece la cuestión filosófica: ¿qué es la probabilidad?

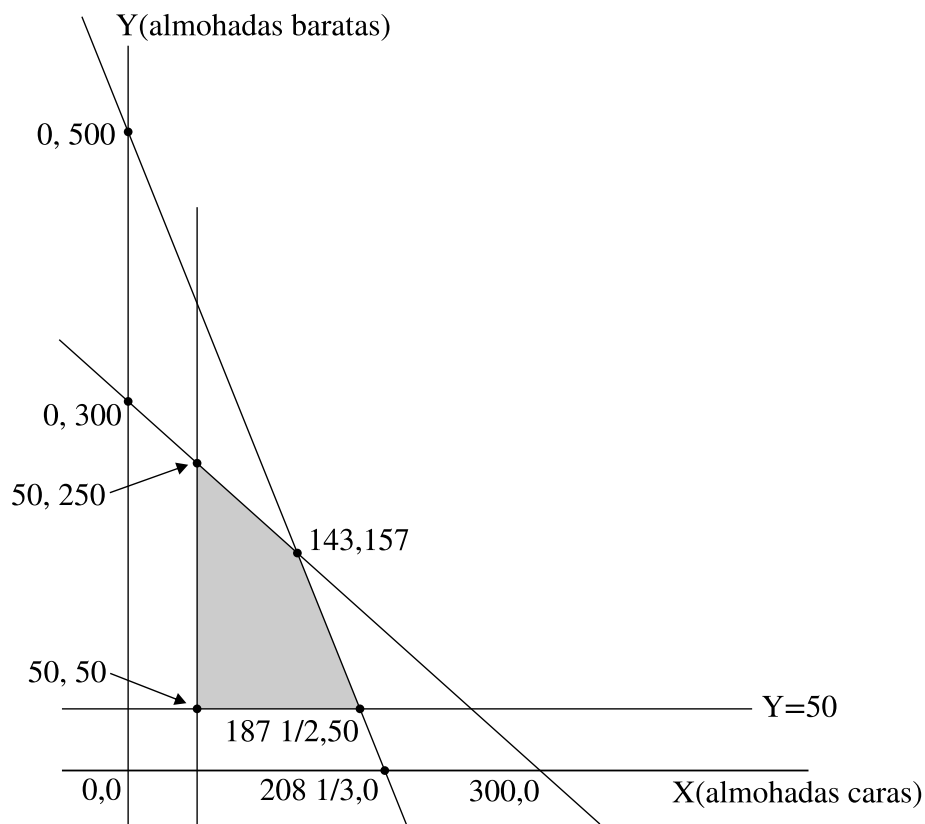
Programación lineal

La programación lineal es un método para maximizar (o minimizar) una cierta cantidad asegurando al mismo tiempo que se cumplen ciertas condiciones sobre otras cantidades. Generalmente estas condiciones son lineales (sus gráficas son líneas rectas), de ahí el nombre de la disciplina: programación lineal. Es una de las técnicas más útiles de la investigación operacional, que es como se conoce el conjunto de instrumentos matemáticos desarrollados después de la segunda guerra mundial para mejorar el rendimiento de los sistemas económicos, industriales y militares, y desde entonces se ha convertido en un ingrediente habitual de los cursos de matemáticas de las escuelas de empresariales. (Véase la entrada sobre *El método de simulación de Montecarlo y Matrices*).

En vez de seguir invocando inexpresivos términos matemáticos para aclarar su significado, lo ilustraremos reflexionando sobre un simple cálculo del punto muerto. Un pequeño taller fabrica sillas metálicas (o artefactos si prefiere las formulaciones genéricas). Sus costes son 80.000 ptas. (en bienes de equipo, por ejemplo) y 3.000 ptas. por cada silla producida. Así pues, el coste total T contraído por el taller viene dado por la fórmula $T = 3.000X + 80.000$, donde X es el número de sillas producidas. Si suponemos además que el precio de venta de estas sillas es de 5.000 ptas. la pieza, los ingresos totales R del taller vienen dados por la ecuación $R = 5.000X$, donde X es el número de sillas vendidas.

Representando ambas ecuaciones sobre el mismo par de ejes coordenados, encontramos que se cortan en un punto en el cual los costes y los ingresos son iguales. El punto muerto, o de beneficio cero, es el (40, 200.000 ptas.), de modo que si se venden menos de 40 sillas, los costes

superan los ingresos; si se venden más, los ingresos superan los costes; y si se venden exactamente 40 sillas, tanto los ingresos como los costes son 200.000 ptas. Maximizar los beneficios en este caso se reduce a vender tantas sillas como sea posible. (Para obtener algebraicamente el punto de beneficio cero, 40, se resta la ecuación $Y = 3.000X + 80.000$ de la $Y = 5.000X$. La ecuación resultante, $0 = 2.000X - 80.000$, se resuelve fácilmente y da $X = 40$).



$$X = 50 \quad 12X = 5Y = 2500 \quad X + Y = 300$$

La región sombreada satisface todas las desigualdades

Después de este preliminar, consideremos el siguiente problema, que es un caso auténtico de programación lineal. Sin dejar las aplicaciones de la economía, supondremos que una empresa fabrica dos tipos de almohadas. Producir una almohada cara cuesta 1.200 ptas. y se vende a 3.000 ptas., mientras que una barata cuesta 500 ptas. y se vende a 1.800 ptas. La compañía no puede fabricar más de 300 almohadas al mes y no puede gastar

más de 250.000 ptas. al mes en su producción (son normas impuestas por la subvención).^[15] Si la compañía ha de fabricar al menos 50 almohadas de cada tipo ¿cuántas ha de fabricar de cada clase para maximizar sus beneficios?

Si llamamos X al número de almohadas caras que la compañía fabrica cada mes e Y al de almohadas baratas, podemos convertir las condiciones sobre X e Y del problema en: $X + Y \leq 300$; $X \geq 50$; $Y \geq 50$; y $1.200X + 500Y \leq 250.000$. La última desigualdad se debe a que si fabricar una almohada cara cuesta 1.200 ptas., producir X costará $1.200X$ ptas.; y análogamente, hacer Y almohadas baratas costará $500Y$ ptas. Obsérvese que estas condiciones se expresan como desigualdades lineales, cuyas gráficas son regiones del plano delimitadas por líneas rectas (o, en problemas más complicados, por sus análogos en espacios de más dimensiones).

La cantidad que hay que maximizar es el beneficio, que en términos de X e Y vale $P = 1.800X + 1.300Y$. Esto es así porque el beneficio que se tiene por cada almohada cara es de 1.800 ptas. (3.000 ptas. – 1.200 ptas.), y por cada almohada barata 1.300 ptas. (1.800 ptas. – 500 ptas.), con lo que X de las primeras dan un beneficio de $1.800X$ ptas., e Y de las segundas dan $1.300Y$ ptas. Una vez tenemos el problema planteado así, hay varias técnicas para hallar la solución. Una es gráfica y consiste en encontrar los vértices y los lados de la región permitida —la parte del plano en la que son válidas todas las desigualdades— y luego probarlas para encontrar en cuál de ellas se tiene el máximo beneficio. Con este método, y un poco de geometría analítica, descubrimos que la compañía de almohadas debería fabricar 143 almohadas caras y 157 baratas al mes si quiere obtener el máximo beneficio.

Otra técnica, llamada método simplex, debida al matemático norteamericano George Danzig, desarrolla y formaliza esta estrategia geométrica de modo que un ordenador pueda examinar rápidamente estos puntos en el caso de que haya más de dos variables. Usado durante más de cuarenta años, el método simplex ha ahorrado una cantidad incalculable de tiempo y dinero. Sin embargo, si el problema de optimización tiene varios miles de variables y desigualdades lineales, como ocurre por ejemplo al establecer el horario de unas líneas aéreas o los recorridos de las llamadas telefónicas, la comprobación puede ser un poco lenta, incluso para un ordenador. Para estas ocasiones existe un algoritmo, inventado recientemente

por Narendra Karmarkar, investigador de los AT&T Bell Laboratories, que a menudo es más rápido en la determinación del horario más eficaz o el recorrido más corto.

Cuando las condiciones no son lineales, los problemas son mucho más difíciles de tratar. Me es grato informarles de que los problemas de programación no lineal frecuentemente colapsan los superordenadores más potentes.

QED, demostraciones y teoremas

Un teorema es una proposición que se deduce, por aplicación únicamente de la lógica, a partir de los axiomas aceptados y de otras proposiciones demostradas previamente. Normalmente, sólo se da el nombre honorífico de teorema a proposiciones y enunciados que son importantes y principales. Una consecuencia inmediata de un teorema se llama corolario de dicho teorema, mientras que un lema es un enunciado, habitualmente técnico, que se precisa para demostrar el teorema. Los dibujos, diagramas y ejemplos pueden hacer que un enunciado sea creíble, pero lo único que lo convierte en teorema es una demostración detallada.

Naturalmente ésta no es más que la historia oficial. El autor de un teorema de una revista de investigación matemática, especialista en, digamos grupos de hemi-semi-demioperadores de orden exponencial primo, normalmente esboza unos argumentos que le convencen a él, a un par de otros expertos en hemi-semi-demi y al editor. El resultado (los matemáticos suelen llamar «resultados» a los teoremas) es muy probablemente válido, pero usted quizá no pondría la mano en el fuego por él.

Tiene relación con esto una experiencia que he conocido en varios seminarios, coloquios y conferencias. El orador ha llenado la pizarra, o las transparencias, de una densa cortina de definiciones, ecuaciones y demostraciones. Me he perdido, pero me percato de que un buen número de oyentes está asintiendo sabiamente con la cabeza. En una interrupción de la charla, mientras el orador borra la pizarra u ordena sus transparencias, pregunto a uno de los que asentían con entusiasmo sentado a mi lado qué significa uno de los símbolos cruciales. Por su tímido encogimiento de hombros me queda claro que anda tan perdido como yo. La conferencia

continúa y él sigue con su cabeceo afirmativo. Observo que además de los que asienten están los que disienten. Hay también, supongo, unos cuantos matemáticos cuyas especialidades son lo suficientemente afines a la del conferenciante para que no sientan la necesidad de asentir ni la tentación de disentir. Son los guardianes provisionales de la virtud matemática.

En cualquier caso, tradicionalmente se escribían las letras QED al final de la demostración de un enunciado con objeto de subrayar que éste había alcanzado el rango superior de la teoremidad. Son las siglas de la frase latina «*Quod erat demonstrandum*», que significa «Lo que había que demostrar», aunque a veces sirven también para otro fin: la intimidación. Para reaccionar interrogativamente a estas tres letras, que se imprimen en mayúscula y se pronuncian con una inflexión de la voz, hace falta mucha confianza en uno mismo. Demasiada para las posibilidades de la mayoría, especialmente si se tiene en cuenta que la falta de confianza matemática es una condición bastante general en todas las edades. (De hecho no hace falta ser anumérico ni matemáticamente ignorante para que a uno le intimiden así. La frase «Es banal» con la que un matemático eminente despacha la demostración inexistente de un teorema tiene a menudo el mismo efecto intimidador sobre los estudiantes de doctorado que sobre los matemáticos profesionales).

Actualmente, la mayoría de textos acostumbran a indicar que una demostración se ha acabado con una señal vertical en negro, ■, más funcional y menos pretenciosa. Esta práctica fue introducida por el matemático norteamericano Paul Halmos y es preferible, pues sirve igual que el QED para indicar el final, sin esa connotación intimidadora. ¿Se puede acaso pronunciar ■ con una inflexión en la voz? No obstante, yo sostengo que no habría que renunciar al empleo del QED en ocasiones señaladas, como los principales teoremas, pues la locución confiere al demostrador una sensación más solemne de satisfacción y finalidad que el un poco plebeyo ■. Al fin y al cabo, hay un límite a lo que uno puede hacer para evitar la intimidación.

La lógica matemática ha cambiado una barbaridad en los últimos 2.500 años. Los silogismos de Aristóteles llevaron a las clasificaciones medievales de los razonamientos, que a su vez llevaron a las álgebras de Boole de proposiciones. Los lógicos de finales del siglo XIX y del XX, como Frege, Peano, Hilbert, Russell y Gödel, han rigorizado y generalizado enormemente

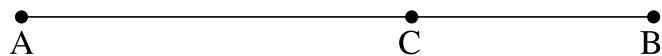
las lógicas clásica y medieval y han creado el potente aparato de la moderna lógica de predicados. Sin embargo, la esencia de la lógica y el atractivo cautivador de la demostración matemática siguen ahí y se reflejan en las tres letras QED. Significan, abreviadísimamente, que el teorema se sigue *necesariamente* de las hipótesis y que (si se ha hecho bien) nada ni nadie pueden cambiarlo.

Una última cosa sobre las demostraciones. Mucha gente piensa que sólo son aceptables aquellas demostraciones que se expresan en forma simbólica y utilizan toda la parafernalia de la lógica formal. Sin embargo, las más de las veces tales demostraciones no hacen sino embrollar las cosas. Con mucho es preferible un razonamiento verbal claro y convincente. Véase, por ejemplo, la demostración de que 6 es el menor número de invitados necesarios que garantizan que al menos 3 de ellos se conocen o que 3 de ellos no se conocen en la entrada sobre *Combinatoria*, o la de la propiedad del punto fijo del escalador en la entrada sobre *Topología*.

Rectángulo áureo, sucesiones de Fibonacci

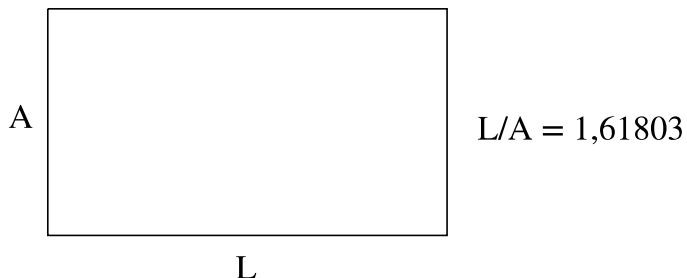
Si un matemático de la Grecia antigua viajara en el tiempo y aterrizara en una moderna tienda de artículos de oficina, una de las muchas cosas que quizá le impresionaran serían las fichas de 3 por 5 pulgadas. Después de admirar por un instante las reglas y compases y luego las carpetas de piel y las fantásticas calculadoras, probablemente volvería a las fichas de nuevo, esta vez fascinado por las de 5 por 8 pulgadas. La razón de su fascinación por las fichas (ya sé que es perfectamente posible que no le interesaran lo más mínimo, pero supongamos que sí) consistiría en que sus dimensiones son aproximadamente las del rectángulo áureo, una figura que nuestro matemático griego y sus contemporáneos habían considerado muy interesante.

Antes de volver a las fichas, definiré primero el concepto de sección áurea, una razón que para muchos resulta muy armoniosa y que está íntimamente relacionada con el tema que nos ocupa. Imaginemos que tenemos un segmento de recta AB y queremos dividirlo por algún punto interior C. Podríamos escoger C de modo que fuera el punto medio entre A y B, pero ésta sería una elección muy insulsa, de manera que supondremos que C divide AB en una parte más larga AC y otra más corta CB. Los pitagóricos nos aconsejarían que escogiéramos C de modo que la razón entre el segmento entero y la parte larga sea igual a la razón entre la parte larga y la corta, esto es, $AB/AC = AC/CB$. Si tomamos C de este modo, se dice que C divide AB en su sección áurea, y se puede calcular que esta razón áurea (entre el todo y la parte larga o entre la parte larga y la corta) es aproximadamente de 1,61803 a 1 (excepto quizás en Wall Street, donde es muy probable que la razón áurea sea precio/ganancias).



Se dice que C divide AB en la proporción áurea si $AB/AC=AC/CB$

Las fichas de 3x5 son una aproximación, pues $5/3 = 1,6$



Un rectángulo cuyas dimensiones de longitud y anchura guarden la proporción áurea se llama rectángulo áureo

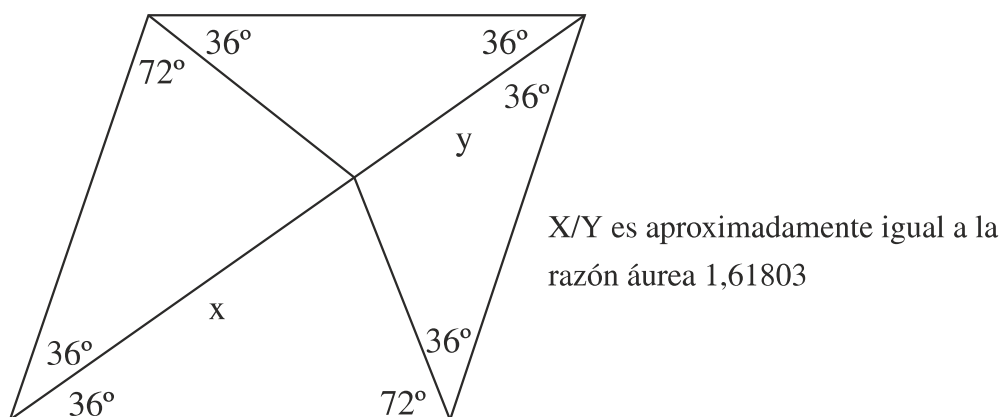
Un rectángulo áureo se define como cualquier rectángulo tal que el cociente entre su longitud y su anchura sea igual a la razón áurea. No es sorprendente pues que el Partenón de Atenas pueda enmarcarse en un rectángulo áureo, al igual que muchas de sus partes. Muchas otras obras de arte griegas utilizaron las proporciones del rectángulo áureo, como también lo hicieron artistas posteriores desde Leonardo a Mondrian y Le Corbusier. La famosa sucesión de Fibonacci 1, 1, 2, 3, 5, 8, 13, 21, 34, 55, 89, 144, 233, ... guarda una relación inesperada con los rectángulos áureos y los relaciona con las fichas de las que hablamos al principio. La sucesión se define por la propiedad de que cada término (excepto los dos primeros) es igual a la suma de los dos que le preceden: $2 = 1 + 1$; $3 = 2 + 1$; $5 = 3 + 2$; $8 = 5 + 3$; $13 = 8 + 5$. (Véase también la entrada sobre *La recurrencia*).

Recordando que la razón áurea es aproximadamente 1,61803 (el número tiene una expresión decimal infinita y no periódica) y haciendo algunas divisiones, vemos (se puede demostrar con un poco de álgebra) que el cociente de un término de la sucesión de Fibonacci y su antecesor tiende a esta razón. Para las fichas de 3 por 5, el cociente $5/3 = 1,66666$; para las de 5 por 8 el cociente es $8/5 = 1,6$; $13/8 = 1,625$; $21/13 = 1,615384$; $34/21 = 1,61905$; etc. Si los griegos estuvieran en lo cierto, los blocs de 8 por 13 pulgadas quizá se venderían más que los de 8 1/2 por 11 pulgadas.

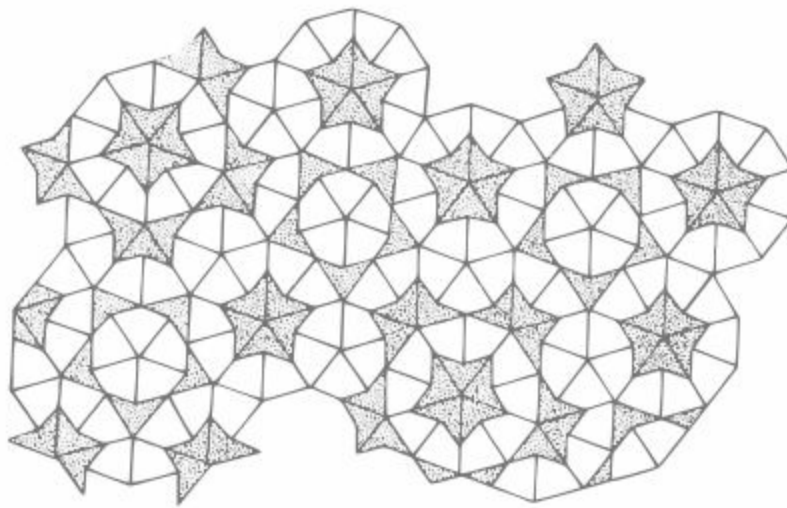
La sucesión de Fibonacci también está presente en otros lugares además de las tiendas de material de oficina. En los girasoles, por ejemplo, el número de espirales hacia la izquierda y el número de espirales hacia la derecha son generalmente números de Fibonacci sucesivos. Del mismo modo, el número de conejos en generaciones sucesivas parece seguir una pauta de Fibonacci, y la concha del *nautilus* se puede generar fibonáccicamente (por inventar un adverbio horrible).

El rectángulo áureo y la armonía estática que representa es típico de la geometría griega clásica, mientras que la sucesión de Fibonacci, que data de cerca del año 1200, sugiere el lento despertar de una concepción más cuantitativa y numérica de la matemática. Uno y otra evocan una placidez que parece un tanto disonante con nuestra era actual, más fracturada y espinosa, cuyo emblema matemático más apropiado es la teoría del caos.

Pero la matemática no respeta en lo más mínimo las pomposas declaraciones históricas y, a principios de los años setenta, el físico matemático inglés Roger Penrose descubrió un nuevo ejemplo de razón áurea con un sabor algo más moderno. Halló dos figuras sencillas (una en forma de cometa y la otra en forma de dardo) tales que con réplicas de ellas se puede recubrir el plano de una manera no periódica y cuyas dimensiones guardan la proporción áurea. Además, y éste es el lado moderno, con ellas no puede recubrirse el plano periódicamente.



Las dos figuras simples de Penrose (un dardo y una cometa) y cómo encajan una en otra



Parte de un recubrimiento no periódico del plano con las dos figuras simples de Penrose

La recurrencia: de las definiciones a la vida

La expresión $5!$ indica el producto $5 \times 4 \times 3 \times 2 \times 1$, $19!$ significa el producto $19 \times 18 \times 17 \times \dots \times 3 \times 2 \times 1$. Estas locuciones se leen «5 factorial» y «19 factorial», respectivamente, y no «¡cinco!» ni «¡diecinueve!» con una entonación exclamativa. Aunque principalmente se usan en probabilidad y en otras ramas de la matemática en las que hace falta contar todas las posibles realizaciones de un suceso, a veces pienso que sería bueno que la gente considerara a los demás como «factoriales históricos». Así, (Marta)! no significaría la Marta actual, sino el producto de todas sus experiencias pasadas.

Formalmente, establecemos que $1!$ es 1, y luego definimos $(N + 1)!$ como $(N + 1) \times N!$. Es decir, definimos el «factorial» explícitamente para el primer término y luego definimos su valor para cualquier otro término mediante los valores de los términos anteriores al mismo. Una definición de este tipo se llama recurrente, y podemos usarla para calcular el valor de $5!$. Vale 5 por $4!$. Pero ¿cuánto vale $4!$? Pues, 4 por $3!$. ¿Y $3!$? 3 por $2!$. ¿Y cuánto es $2!$? Pues, 2 por $1!$. Y, por fin, la definición nos dice cuánto vale $1!$. Es 1. Reuniéndolo todo encontramos que $5! = 5 \times 4 \times 3 \times 2 \times 1$. Tranquilos, no repetiré esta letanía para $19!$.

De modo análogo, la definición recurrente de la suma nos dice cómo sumar cualquier número a X . Establece primero que $X + 0$ es X y define por recurrencia que $X + (Y + 1)$ es igual a 1 más $X + Y$. (Ruego me disculpen si el resto del párrafo parece un chiste de esos que se estiran como una goma). Para determinar $(8 + 3)$ aplicando la definición anterior, por ejemplo, observaremos que $8 + 3 = (8 + 2) + 1$, que $8 + 2 = (8 + 1) + 1$, que $8 + 1 = (8 + 0) + 1$ y que $8 + 0 = 8$. Reuniéndolo todo tenemos que $8 + 3 = 8 + 1 +$

1 + 1. Y así hemos reducido la suma a contar. La definición por recurrencia de la multiplicación establece que $X \times 0 = 0$ y luego se define $X \times (Y + 1)$ igual a $(X \times Y)$ más X . Aplicando la definición podemos determinar el valor de 23×9 reduciéndolo a una serie de sumas [$23 \times 9 = (23 \times 8) + 23$; $23 \times 8 = (23 \times 7) + 23$; $23 \times 7 = (23 \times 6) + 23$; ...] las cuales se reducen a su vez a contar. Análogamente podemos definir el significado de elevar X a una potencia cualquiera. Primero establecemos que X^0 es 1 y luego definimos $X^{(Y+1)}$ como X por X^Y . Por tanto, $7^4 = 7 \times 7^3$, $7^3 = 7 \times 7^2$, etc. la exponenciación se ha reducido a la multiplicación, la cual se reduce a la suma, que a su vez se reduce a contar.

Aunque de entrada pueda parecer una tontería, la idea de expresar el valor de una función de $(N + 1)$ por recurrencia, en términos de sus valores anteriores, es muy útil y en informática es inevitable. De hecho, la recurrencia constituye el mismísimo núcleo de la programación informática, con su empleo característico de bucles (la realización repetitiva de un mismo procedimiento para distintos valores de una misma variable), subrutinas y otras estrategias que permiten reducir procesos complejos a simples operaciones aritméticas. Las funciones matemáticas y los algoritmos que admiten una definición por recurrencia son precisamente los únicos que se pueden manejar con un ordenador. Es decir, si una función es recurrente, puede calcularla un ordenador, y si un ordenador puede calcular cierta función, ésta es recurrente. Además, estas definiciones recurrentes se pueden encajar unas en otras, se pueden iterar indefinidamente y, mediante las codificaciones y correspondencias apropiadas, se pueden generalizar para toda clase de actividades, aunque éstas no tengan aparentemente mucho que ver con el cálculo.

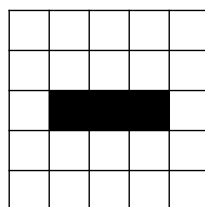
Estas funciones y definiciones tienen un papel importante en lógica pues precisan lo que quiere decirse con palabras tales como «mecánico», «regla», «algoritmo» y «demostración» (véase también la entrada sobre *La inducción matemática*). En gramática formal los lingüistas usan definiciones recurrentes para clarificar las reglas de la gramática y estudiar los procesos cognitivos. Demuestran cómo se pueden construir largos enunciados complejos a partir de frases y oraciones cortas. Si se usa en combinación con la autorreferencia,

la recurrencia es todavía más potente. Algunos virus informáticos, por ejemplo, se reproducen de un modo parecido al de la frase siguiente que proporciona ella misma las instrucciones y el material para su propia réplica. *Alphabetize and append, copied in quotes, these words: «these append, in Alphabetize and word: quotes, copied»*. Dicho en español y sin tanta concisión, la frase anterior manda ordenar alfabéticamente las palabras que siguen a los dos puntos y luego añadir, entre comillas y sin ordenar, estas mismas palabras. Y ¡abracadabra!: la frase se ha replicado a sí misma y lo mismo harán sus descendientes.

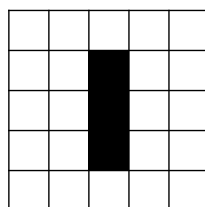
	1	2	3	
	8		4	
	7	6	5	

Un cuadrado marcado
y sus 8 vecinos

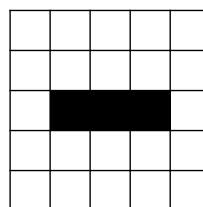
Paso de una generación a la siguiente: los cuadrados marcados con 2 o 3 vecinos marcados siguen. Aquéllos con 0, 1 o más de 4 vecinos marcados desaparecen. Un cuadrado sin marcar que tenga exactamente 3 vecinos marcados resucita



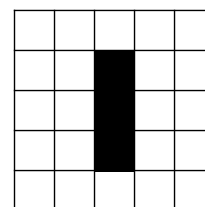
Inicio



Tiempo 1



Tiempo 2



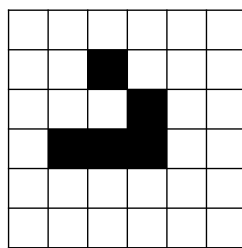
Tiempo 3

Estructura oscilante

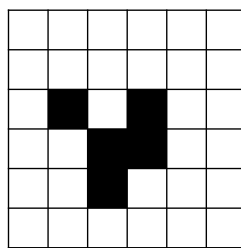
La recurrencia tiene también un papel cada vez más importante en la descripción de fenómenos físicos, especialmente después de los avances de la teoría del caos, que ilustran gráficamente lo naturales y complicadas que pueden llegar a ser las estructuras definidas por recurrencia.

Relacionado con esto último, consideremos una sugestiva aplicación de una definición recurrente simple. El juego solitario «Vida» del matemático británico John Conway transcurre en un tablero de damas infinito que tiene alguno de sus cuadros ocupados por fichas. (El papel cuadriculado con

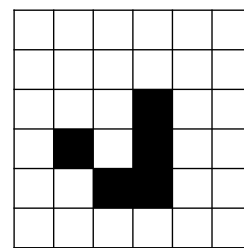
algunos cuadros en oscuro también vale). Como cada cuadro del tablero tiene 8 vecinos (4 adyacentes y 4 en diagonal), cada ficha puede tener entre 0 y 8 vecinas. La distribución de fichas originaria es la primera generación, y el paso de una generación a la siguiente se rige por tres reglas. Cada ficha con 2 o 3 vecinas permanece en el tablero y pasa a la siguiente generación. Cada ficha con 4 vecinas o más se saca del tablero y no pasa a la generación siguiente, y lo mismo ocurre con las que tienen 0 o 1 vecina. Cada cuadro vacío que tenga exactamente 3 vecinos ocupados estará también ocupado por una ficha en la siguiente generación.



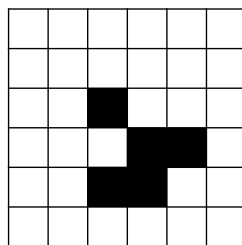
Inicio



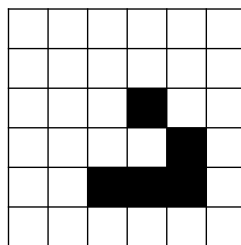
Tiempo 1



Tiempo 2



Tiempo 3



Tiempo 4

Estructura que se reproduce a sí misma

Los cambios dictados por las tres reglas tienen lugar todos a la vez, y el tictac discreto del reloj marca el paso de una generación a la siguiente para este «autómata celular». Es sorprendente cómo estas simples reglas recurrentes que rigen la supervivencia, muerte y nacimiento de las fichas pueden conducir a estructuras y movimientos sobre el tablero de una gran belleza y complejidad. Figuras parecidas a trenes y aviones saltan al tablero a medida que se van sucediendo las distintas generaciones. Algunas configuraciones iniciales mueren, otras se reproducen, mientras que otras

engendran universos enteros de sapos, barcos y dibujos geométricos. Algunas oscilan o parpadean, otras sufren una sucesión cíclica de transformaciones mientras que otras nunca dejan de evolucionar.

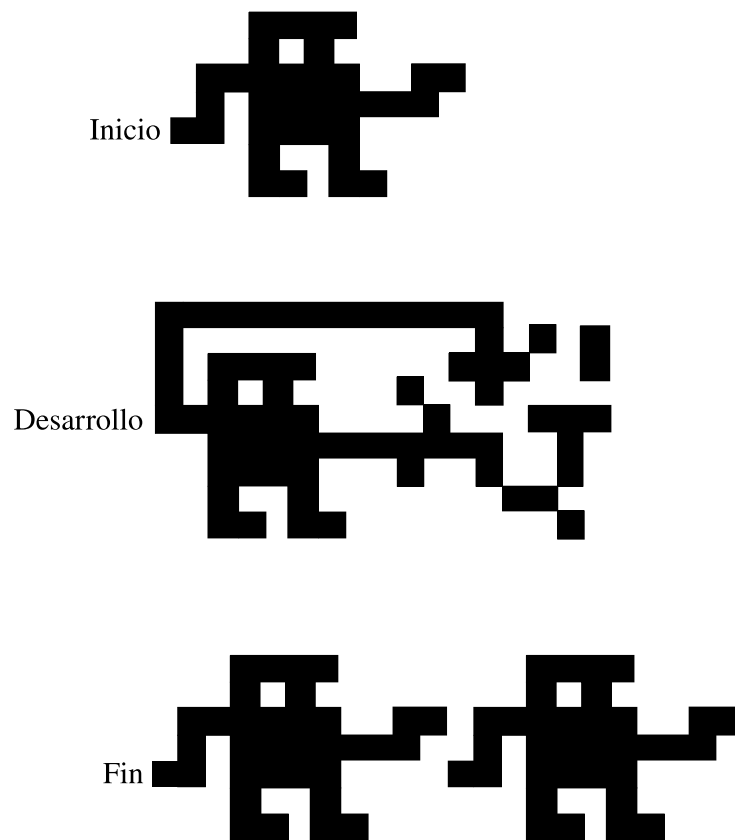


Diagrama simplificado de una figura que se reproduce

Para hacerse una idea de cómo son estas estructuras cambiantes hay que empezar, ya sea con algunas fichas sobre un tablero de damas ya sea con cuadros oscuros sobre papel cuadriculado, y seguir las reglas sistemáticamente, o con una versión del juego para ordenador. Pruebe con varias distribuciones iniciales de fichas. Podría empezar, por ejemplo, con tres fichas en fila. Se obtiene una figura parpadeante, que oscila entre la vertical y la horizontal. Pruebe con algunas más. La configuración inicial de 4 fichas, 3 dispuestas en fila horizontal y 1 sobre la ficha del medio, da lugar a una evolución particularmente interesante.

Conway ha llegado a demostrar que ciertas configuraciones iniciales de fichas se pueden usar como ordenador, lento pero hecho y derecho, cosa que

nos hace volver a las funciones recurrentes en los enteros. La esencia de «Vida», y posiblemente también de la vida, es la recurrencia (véase la entrada sobre *Test de Turing*). Los matemáticos, desde Pitágoras a Poincaré, han escrito que los números y el contar son la base de toda la matemática. Las definiciones recurrentes nos dan un punto de vista desde el que esta pretensión es verdaderamente modesta.

La regla del producto

La rima de la «St. Ivés Mother Goose» [«La mamá oca de St. Ivés»] plantea un problema que ya aparece en una forma casi idéntica en el papiro Rhind del antiguo Egipto que data del año 1650 a. C.: «Yendo hacia St. Ivés me encontré un hombre con siete esposas. Cada esposa tenía siete sacos y cada saco contenía siete gatos, cada gato tenía siete gatitos. Gatitos, gatos, sacos, y esposas ¿cuántos iban a St. Ivés?». La respuesta es uno, pues todos los demás iban en dirección contraria a St. Ivés, pero la determinación del tamaño de un grupo depende de la comprensión de la regla del producto.

La matemática combinatoria no contiene ninguna otra idea tan simple, y a pesar de ello tan potente, como esta inofensiva regla: si se puede realizar tal acción o hacer una elección en M modos diferentes y luego se puede realizar otra acción u otra elección en N modos distintos, entonces se pueden realizar esas dos acciones o esas dos elecciones consecutivamente en $M \times N$ modos distintos.

Tomando un ejemplo menos antiguo, suponga que se encuentra en Los Angeles y ha de llegar como sea a la Costa Este de Estados Unidos; cualquier lugar de la Costa Este sirve. Consulta los horarios de las líneas aéreas y descubre que no hay vuelos directos, pero sí los hay a tres ciudades del Medio Oeste (Minneapolis, Chicago y San Luis, pongamos por caso), y de cada una de ellas parten vuelos a cuatro ciudades de la Costa Este (Boston, Nueva York, Filadelfia y Washington, por ejemplo). La regla del producto nos dice que hay 12 ($= 3 \times 4$) maneras distintas de llegar a la Costa Atlántica. Se pueden simbolizar estas 12 maneras por: MB, MN, MF, MW, CB, CN, CF, CW, SB, SN, SF y SW, donde las siglas corresponden a las iniciales de las ciudades citadas.

La regla puede aplicarse más de una vez. Por ejemplo, si el alumnado de una escuela consta de 18.000 alumnos, podemos estar absolutamente seguros de que al menos dos de ellos tienen las mismas tres iniciales. La razón es que las 26 primeras iniciales posibles pueden ir seguidas de cualquier otra de las 26 iniciales centrales posibles, las cuales pueden a su vez ir seguidas de otra cualquiera de las 26 iniciales últimas posibles. Así pues, por la regla del producto, hay 26^3 o 17.576 conjuntos de tres iniciales (ordenadas). Como este número es menor que el de alumnos matriculados podemos concluir que por lo menos dos estudiantes han de tener las mismas iniciales. (De hecho, esto ocurre al menos con 424 alumnos y, en términos prácticos, es probable que sean muchos más).

El alfabeto Braille, cuyos símbolos consisten en combinaciones de dos columnas verticales de tres puntos cada una, nos proporciona otro ejemplo. Las distintas letras y símbolos de este alfabeto para ciegos se distinguen por el subconjunto de los seis puntos que están en relieve. Así por ejemplo, la letra «a» se indica poniendo en relieve sólo el punto superior del lado izquierdo, mientras que la letra «r» se indica dejando en relieve los tres puntos de la columna izquierda y el punto central de la derecha. ¿Cuántos símbolos hay en total? Tenemos dos posibilidades para cada punto: en relieve o no. Por tanto, como hay seis puntos, tenemos 2^6 posibilidades distintas. Como una de ellas no tiene ningún punto en relieve, es imperceptible, por lo que hay 63 símbolos Braille distintos (letras, números, combinaciones de letras, palabras más corrientes y signos de puntuación).

O considérese un mensaje codificado en inglés que haya de tener la forma SPOOK7, de modo que los dos primeros signos sean consonantes, los dos siguientes vocales, el siguiente una consonante y el último un número comprendido entre 1 y 9. Hay $(21^2 \times 5^2 \times 21 \times 9)$ o 2.083.725 posibles mensajes de este tipo. Si todos los signos del mensaje han de ser distintos sólo hay $(21 \times 20 \times 5 \times 4 \times 19 \times 9)$ o 1.436.400 mensajes de esta segunda clase. El número de teléfonos posibles en una provincia es aproximadamente 8×10^6 , pues el primer número puede ser cualquier dígito distinto de 0 o 1 y los seis restantes, cualquiera de los 10 dígitos. (En la práctica los números telefónicos están sujetos a más condiciones que no hemos tenido en cuenta

aquí).

Si las placas de matrícula de una cierta provincia siguen la pauta NNNN-LL, cuatro dígitos seguidos de dos letras, el número de placas de matrícula posibles de dicha provincia es de $(10^4 \times 26^2)$ o de 6.760.000. Si las letras y los dígitos hubieran de ser todos distintos, tendríamos sólo $(10 \times 9 \times 8 \times 7 \times 26 \times 25)$ o 3.276.000 placas distintas.

Los números que usan los grandes almacenes, las empresas de servicios públicos y las compañías de tarjetas de crédito para identificarnos suelen constar de 15 o 20 símbolos, muchísimo más de lo que haría falta, atendiendo a la población de Estados Unidos. Aun consistiendo sólo en dígitos, una sucesión de 20 símbolos basta para asignar un número de identificación a 10^{20} personas, 20.000 millones de veces toda la población mundial. Una utilidad de esta capacidad extra consiste en hacer muy improbable que un posible impostor o un estafador pueda dar por casualidad con una sucesión que corresponda a un cliente real.

Como ilustran estos ejemplos, una consecuencia sorprendente de la regla del producto es el gran número de posibilidades que resultan de su aplicación repetida. Este número crece exponencialmente con el número de veces que se aplica la regla, y hasta los ordenadores más rápidos, al intentar enumerar todas las posibilidades o al aplicar métodos de fuerza bruta para resolver problemas complejos, enseguida tropiezan con la dificultad conocida como explosión combinatoria y van a paso de tortuga.

El mismo rebrotar de posibilidades (aunque a una escala más modesta) es un problema que fastidia los intentos de algunos autores o algunos directores de escribir un libro o hacer una película en la que haya un cierto número de coyunturas en las que el lector o el espectador puedan expresar su deseo y elegir el modo de continuar. Con sólo 5 de tales situaciones habría que escribir 32 libros distintos o hacer 32 películas para acomodarse a todas las posibles elecciones. (Habría dos opciones en la primera coyuntura, cada una de las cuales llevaría a una nueva coyuntura, que tendría a su vez dos opciones que llevarían a una coyuntura cada una, etc., de lo que resultarían en total 2^5 o 32 tratamientos distintos). Si hubiera más coyunturas o más opciones en cada una de ellas, el número sería muchísimo mayor. De hecho,

el número de obras necesarias que un autor o director habría de realizar para simular la sensación de libertad inherente incluso a una conversación breve con alguien, le obligaría a dedicar toda la vida a explorar todas las posibles representaciones. (Véase la entrada sobre *La conciencia humana*).

Esta idea de ramificación continua subyace en las muchas metáforas acerca de amigos que se separan, historias divergentes y personas que se vuelven excéntricas con los años. También juega un papel importante, como ya hemos apuntado, en nuestra concepción de la libertad y (para poner un ejemplo más extravagante) en la llamada interpretación de los universos paralelos de la mecánica cuántica, en la que el universo se descompone a cada instante en una infinidad de universos incomunicados.

La regla del producto tiene muchas otras aplicaciones y variantes. De entre ellas, las más útiles y mejor conocidas hacen intervenir los números combinatorios. (Véase la entrada acerca del *Triángulo de Pascal* para una discusión acerca de los mismos).

[El número del grupo del poema del St. Ivés Mother Goose es 2.801. ¿Cuál sería la respuesta si cada gatito tuviera siete pulgas?].

Series: convergencia y divergencia

Las series infinitas y sus aplicaciones constituyen un importante dominio del análisis matemático. Las series fueron usadas informalmente por los matemáticos mucho antes de que se las llegara a comprender plenamente (véase la entrada sobre *Zenón*), y todavía siguen resultando seductoras para nuestra intuición de los conceptos de número e infinito. Hablando un poco vagamente (N simbolizará un entero positivo arbitrario y... indicará que la serie sigue indefinidamente), afirmo que la suma $1 + 1/3 + 1/9 + 1/27 + 1/81 + 1/243 + 1/729 + \dots + 1/3^N + \dots$ es finita, mientras que la suma $1 + 1/2 + 1/3 + 1/4 + 1/5 + 1/6 + \dots + 1/N + \dots$ es infinita. ¿Por qué esta diferencia? Y también, bueno, ¿y qué?

A grandes rasgos, la respuesta a la primera pregunta es que los términos de la primera serie disminuyen lo suficientemente deprisa para que la suma sea acotada, mientras que no ocurre lo mismo con los de la segunda serie. En la primera serie se pasa de un término al siguiente dividiendo por 3. En la segunda los términos se van reduciendo muy lentamente y puede demostrarse que si se suma un número suficientemente grande de ellos se obtiene un resultado mayor que cualquier número fijado de antemano —mil trillones, por ejemplo. Se dice que la primera serie converge (su suma es $3/2$ o $1\ 1/2$) y que la segunda diverge—. (Las palabras «converge» o «diverge» pueden sonar un tanto raras aquí pero son las que se usan corrientemente en matemáticas).

No siempre es evidente si una serie converge o diverge. La serie $1 + 1/4 + 1/9 + 1/16 + 1/25 + \dots + 1/N^2 + \dots$ converge (maravillosamente, su suma es $\pi^2/6$), mientras que $1/\log(2) + 1/\log(3) + 1/\log(4) + 1/\log(5) + \dots$ diverge.

(Para los aficionados a la notación, señalaré que frecuentemente se usa la letra griega $\sum$ para indicar una serie y ∞ para simbolizar el infinito; así pues $\sum_{N=1}^{\infty} \frac{1}{N^2} = \frac{\pi^2}{6}$ y $\sum_{N=2}^{\infty} \frac{1}{\log(N)}$ diverge). La serie cuyos términos consisten en 1 dividido por los sucesivos factoriales, $1 + 1/1! + 1/2! + 1/3! + 1/4! + 1/5! + \dots + 1/N! + \dots$, converge (también maravillosamente, me parece) al número e , mientras que la serie $1/1.000 + 1/(1.000 \times \sqrt{2}) + 1/(1.000 \times \sqrt{3}) + 1/(1.000 \times \sqrt{4}) + 1/(1.000 \times \sqrt{5}) + \dots + 1/(1.000 \times \sqrt{N}) + \dots$ diverge.

El significado de convergencia de una serie se puede aclarar mejor mediante el concepto de «suma parcial». La idea consiste en llegar a la «suma infinita» por medio de una sucesión de sumas parciales. Para las tres series convergentes citadas esto significa que: la sucesión de sumas parciales $1, 1 + 1/3, 1 + 1/3 + 1/9, 1 + 1/3 + 1/9 + 1/27, \dots$ se aproxima tanto como queramos a $1\frac{1}{2}$; las sumas parciales $1, 1 + 1/4, 1 + 1/4 + 1/9, 1 + 1/4 + 1/9 + 1/16, \dots$ se aproximan a $\pi^2/6$ con una precisión arbitrariamente grande; y análogamente, las diferencias entre las sumas parciales $1, 1 + 1/1!, 1 + 1/1! + 1/2!, 1 + 1/1! + 1/2! + 1/3!, \dots$ y la constante matemática e se hacen tan pequeñas como queramos (véase la entrada sobre E para su definición).

Las series llamadas geométricas son relativamente fáciles de manejar. En ellas, cada término se obtiene multiplicando el anterior por un factor constante. La primera de las series citadas es un ejemplo, al igual que $12 + 12(1/5) + 12(1/25) + 12(1/125) + 12(1/625) + \dots$. Las series geométricas *finitas* aparecen en muchos contextos de la vida cotidiana. Si cada año invertimos 100.000 ptas. al 10% de interés en un fondo de pensiones, al cabo de 18 años tendremos en el fondo $100.000 \text{ ptas.} + 100.000(1,1) \text{ ptas.} + 100.000(1,1)^2 \text{ ptas.} + 100.000(1,1)^3 \text{ ptas.} + 100.000(1,1)^4 \text{ ptas.} + \dots + 100.000(1,1)^{18} \text{ ptas.}$, que da aproximadamente 5.600.000 ptas. El primer término de la serie son las 100.000 ptas. que acabamos de invertir. El segundo término de $100.000(1,1) \text{ ptas.}$, el 110% de 100.000 ptas., es el valor de las 100.000 ptas. que invertimos hace un año. El tercer término es el valor acumulado de las 100.000 ptas. invertidas hace dos años, y así sucesivamente hasta llegar al último término de la serie, $100.000(1,1)^{18} \text{ ptas.}$, que es lo que valen hoy las 100.000 ptas. invertidas hace dieciocho años.

Un ejemplo concreto de serie geométrica infinita lo tenemos en las pensiones vitalicias. ¿Cuánto dinero ha de depositar hoy en una cuenta para que usted, sus herederos, y los herederos de sus herederos puedan sacar 100.000 ptas. anuales por siempre jamás? Si suponemos un rédito constante del 10%, es claro que bastará con 1.000.000 ptas. para producir 100.000 ptas. de interés anual.

Quizá es algo más difícil ver que 1.000.000 ptas. es la suma de la siguiente serie infinita: 100.000 ptas. + 100.000/1,1 ptas. + 100.000/(1,1)² ptas. + 100.000/(1,1)³ ptas. + 100.000/(1,1)⁴ ptas. + ... + 100.000/(1,1)^N ptas. + ... El primer término de la serie son las 100.000 ptas. que sacará usted mañana. El término siguiente son las 100.000 ptas. que sacará el próximo año (hay que *dividir* por 1,1 o el 110 porque su valor actual es mucho menor que 100.000 ptas.). El término siguiente son las 100.000 ptas. que sacará dentro de dos años [divididas por (1,1)² porque su valor actual es todavía más pequeño], Y cada término sucesivo se divide por otro factor 1,1 para reflejar la devaluación correspondiente a un año más. Si esta pensión fuera el premio de una lotería, podría llamarse la lotería del fin de los tiempos: 100.000 ptas. anuales por siempre jamás. Es menos engañoso hablar del 1.000.000 ptas. de su valor actual. Del mismo modo, un millón de pesetas al año por siempre jamás sólo son diez millones de pesetas actuales.

Las series geométricas aparecen también cuando nos interesa saber la cantidad de medicina contenida en la sangre de un paciente que la está tomando en un tratamiento diario a largo plazo. Ello se debe a que hemos de sumar la cantidad de medicina en sangre de la tomada hoy, la que queda de la de ayer (una fracción de la de hoy), la de anteayer (la misma fracción de la de ayer), etc. Por razones parecidas también salen esas series cuando nos interesamos ya sea en el impacto total de una compra de títulos del Estado por parte del Banco de España o en la distancia total recorrida por una pelota que bota. La fórmula de la suma de una serie infinita de la forma $A + AR + AR^2 + AR^3 + AR^4 + AR^5 + \dots + AR^N + \dots$, cuando R es menor que 1 y mayor que -1 , es $A/(1 - R)$. A es el término inicial de la serie y R es el factor constante por el que se multiplica cada término para obtener su sucesor. En la primera de las series geométricas citadas A es 1 y R , 1/3, de

ahí que su suma $1/(1-R)$ sea $1/[1 - (1/3)]$ o $3/2$. ¿Cuánto vale la suma de $12 + 12(1/5) + 12(1/25) + 12(1/125) + 12(1/625) + \dots$?

Desgraciadamente, no todas las series son geométricas y se han inventado reglas y tests sofisticados para determinar si son o no convergentes, para calcular su velocidad de convergencia, para encontrar su suma, etc. La suma de una serie, insisto, se define por consideración de la sucesión de sus sumas parciales. En la serie $1 + 1/4 + 1/27 + 1/256 + \dots + 1/N^N + \dots$, esta sucesión es $1, 1 + 1/4, 1 + 1/4 + 1/27, 1 + 1/4 + 1/27 + 1/256, \dots$. Si esta sucesión tiene límite, como es el caso, se dice que este límite es la suma de la serie. O, dicho de otro modo, si la sucesión de sumas parciales se aproxima tanto como queramos a un cierto número, este número (en este caso algo menor que $3/2$) es la suma de la serie. (Véase la entrada sobre *Límites*).

Estas reglas y tests de convergencia son especialmente valiosos para tratar con series de potencias (o polinomios infinitos). Los polinomios normales (finitos) son expresiones algebraicas tales como $3X - 4X^2 + 11X^3, 7 - 17X^2 + 4,7X^5$, o $2X + 5X^3 - 2,81X^4 + 31X^9$. Como el álgebra y el cálculo con esos polinomios es especialmente fácil, los matemáticos han tratado de aproximar otras funciones comunes mediante funciones polinómicas (series de Taylor). (Véase la entrada sobre *Funciones*). Esto puede hacerse para una clase bastante general de funciones. Puede demostrarse, por ejemplo, que la función trigonométrica $\sin(X)$ se puede representar por la serie de potencias (o polinomio infinito) $X - X^3/3! + X^5/5! - X^7/7! + X^9/9! \dots$ y puede aproximarse por una suma parcial de dicha serie. Es decir, $\sin(X)$ es aproximadamente igual a $X - X^3/3! + X^5/5! - X^7/7!$. Análogamente, la función exponencial e^X puede representarse por la serie $1 + X + X^2/2! + X^3/3! + X^4/4! + \dots + X^N/N! + \dots$ y puede aproximarse por la suma polinómica de algunos de los términos primeros. Así, e^2 vale aproximadamente $1 + 2 + 2^2/2! + 2^3/3! + 2^4/4! + 2^5/5!$.

El cálculo de derivadas e integrales, la resolución de ecuaciones diferenciales y el trabajo con números complejos son tareas que se simplifican muchísimo si tratamos con funciones que, como e^X o $\sin(X)$, se pueden representar por series de potencias. De hecho, la importancia de las

series infinitas para el análisis matemático es difícilmente sobreestimable. Los muchos teoremas sobre esta materia son de una gran belleza dentro de la austera estética típica de la matemática.

[La suma de la serie geométrica $12 + 12(1/5) + 12(1/25) + 12(1/125) + 12(1/625) + \dots$ es 15].

Simetría e invariancia

La simetría y la invariancia no son tanto temas matemáticos como principios orientadores de la estética matemática. Desde la preocupación de los griegos por el equilibrio, la armonía y el orden hasta la insistencia de Einstein en que las leyes de la física habían de ser invariantes para todos los observadores, estas ideas han inspirado buena parte del mejor trabajo en matemática y en física matemática.

La simetría y la invariancia son dos conceptos complementarios. Algo es simétrico en la medida que es invariante bajo (o no cambia por) algún tipo de transformación. Para ilustrarlo, consideremos un círculo. Podemos girarlo, o hacer una reflexión con respecto a uno cualquiera de sus diámetros y conserva la misma circularidad. Su simetría consiste en su invariancia bajo estos cambios. Pero supongamos ahora que lo aplastamos un poco (por ejemplo, dibujando el círculo sobre un trozo de madera y comprimiéndola). Observamos que toma una forma elíptica. Ya no se cumple que dos cualesquiera de sus diagonales sean iguales; unas son más largas y otras más cortas. Esta propiedad del círculo se ha perdido, pero otras no. Por ejemplo, el centro de la figura achatada sigue siendo el punto medio de cualquier diámetro, sea cual sea la longitud de éste. Esta última propiedad es invariante aun bajo una transformación tan severa como las de este tipo y refleja pues una clase de simetría más profunda.

Observaciones como ésta sugirieron al matemático alemán del siglo XIX Félix Klein la idea de que los teoremas que atañen a las figuras geométricas podían clasificarse según siguieran o no siendo válidos cuando las figuras se someten a distintos cambios y transformaciones. Con mayor generalidad, dado un cierto conjunto de transformaciones (movimientos rígidos en el

plano, compresiones uniformes, proyecciones), Klein preguntaba qué propiedades de las figuras permanecen invariantes bajo estas transformaciones. El cuerpo de teoremas relativos a estas propiedades es considerado como la geometría asociada a este conjunto de transformaciones.

La geometría euclídea se puede considerar como el estudio de las propiedades que son invariantes bajo movimientos rígidos: rotaciones, traslaciones y reflexiones. Por contra, se entiende por geometría proyectiva la que se ocupa de una clase menor de propiedades que son invariantes bajo todos los movimientos rígidos *más* las proyecciones. (La proyección de una figura es, más o menos, la sombra que proyecta cuando se la ilumina por detrás. La proyección de un círculo podría ser algún tipo de elipse, por ejemplo). Y la topología es la disciplina que se dedica a la clase aún menor de propiedades que son invariantes bajo las transformaciones anteriores y *además* las torsiones, contracciones y estiramientos más severos.

La longitud y el ángulo son propiedades euclídeas (son conservados por los movimientos rígidos), pero no son invariantes bajo transformaciones proyectivas. La linealidad y la triangularidad son propiedades proyectivas (son conservadas por las transformaciones proyectivas, pues las proyecciones transforman siempre las rectas y los triángulos en otras rectas y otros triángulos), pero no son invariantes bajo transformaciones topológicas. Y la conectividad y el número de agujeros de una figura son propiedades que se mantienen a pesar de las torsiones y los estiramientos.

Esta idea de que invariancias profundas indican simetrías más sutiles es muy potente también en ámbitos distintos del puramente geométrico. Las formas de arte simétricas muchísimo más abstractas que, pongamos, el canon griego o la Alhambra de Granada son (en un sentido pickwickiano al menos) el objeto del arte moderno. La teoría especial de la relatividad de Einstein (estuvo pensando en llamarla teoría de los invariantes) fue fruto de su convicción de que las leyes de la física tenían que ser invariantes bajo un grupo de transformaciones descubierto por el físico holandés H. A. Lorentz.

Una consecuencia social del interés de la matemática por las verdades duraderas y eternas y esta estética de la simetría y la invariancia es que, en su forma más pura, la matemática mantiene necesariamente una cierta reserva hacia el mundo real de la contingencia caprichosa y la idiosincrasia humana.

Esta aversión por lo personal se manifiesta incluso en los libros de matemática aplicada y en los de divulgación. Recuerdo haber recibido una carta de un matemático que me decía que le habían gustado mucho mis libros, pero que no eran matemática porque usaba en ellos la palabra «yo» sin ninguna limitación. Tenía en parte razón, desde luego, pero es triste que tuviera que disociar explícitamente el haber disfrutado de los libros y su devoción por la matemática pura. O al menos así lo pienso yo.

La estrategia del matemático consistente en buscar la simetría y la invariancia no puede fracasar, pues el desorden total a todos los niveles del análisis es una imposibilidad lógica. Sin embargo, el descubrimiento de lo asimétrico, lo variable y lo personal no puede hacer daño a nadie.

Sistemas de votación

¿Cómo toman las decisiones las sociedades democráticas? La respuesta es «votando», pero ¿qué significa esto?, si, como suele ocurrir, hay más de dos opciones posibles. Como a menudo un buen ejemplo ilustrativo vale más que páginas y páginas de explicación rigurosa, supongamos, a modo de ilustración, que hay cinco candidatos a la presidencia de una pequeña organización. Aunque todos los miembros del grupo ordenan los cinco candidatos según sus preferencias, el ganador depende críticamente, como veremos, del sistema de votación empleado.

Ahora hay que concretar los números. Supongamos pues que hay 55 electores y que ordenan los candidatos según sus preferencias con el resultado siguiente:

18 electores ordenan los candidatos así: A, D, E, C, B
12 electores ordenan los candidatos así: B, E, D, C, A
10 electores ordenan los candidatos así: C, B, E, D, A
9 electores ordenan los candidatos así: D, C, E, B, A
4 electores ordenan los candidatos así: E, B, D, C, A
2 electores ordenan los candidatos así: E, C, D, B, A

Los partidarios del candidato A quizá digan que habría que usar el método de pluralidad, por el que gana el candidato votado más veces en primer lugar. Con este método gana A fácilmente.

Los partidarios de B quizá digan que debería hacerse una segunda vuelta entre los dos candidatos más votados. En la segunda vuelta B gana con facilidad a A (18 electores prefieren A a B, pero 37 prefieren B antes que A).

La gente del candidato C han de pensar un poco más para encontrar un método que le dé como vencedor. Sugerirán que eliminemos primero al candidato con menos primeros lugares (en este caso E) y que luego reajustemos los votos para el primer lugar de los que quedan (A tiene todavía 18, B tiene ahora 16, C tiene 12, y D sigue con 9). De los cuatro candidatos que quedan eliminamos el que tenga menos primeros lugares y reajustamos la lista de los restantes candidatos (C queda ahora con 21 votos para el primer lugar). Seguimos con este procedimiento de eliminar el candidato con menos votos de primer lugar. Con este método, C se proclama vencedor.

Ahora el director de campaña del candidato D objeta que habría que prestar más atención a la preferencia media, y no sólo a los primeros. Y razona que si se dan 5 puntos a las primeras preferencias, 4 a las segundas, 3 a las terceras, 2 a las cuartas y 1 punto a las quintas, cada candidato tendrá una puntuación, el llamado escrutinio de Borda, que reflejará su popularidad. Como el escrutinio de Borda de D, 191 puntos, es mayor que el de cualquier otro, D gana con este método.

El candidato E es de un temperamento más viril y responde que sólo deberían tenerse en cuenta las luchas hombre a hombre (o hombre a mujer) y que enfrentado a cualquiera de los otros cuatro candidatos en contiendas de dos personas, siempre sale vencedor. Y sostiene por tanto que merece ser el vencedor global. (Alguien que como E gana a todos los demás candidatos de este modo se llama el vencedor de Condorcet. Pero frecuentemente las votaciones son tan embrolladas que no hay ningún vencedor de Condorcet).

¿Quién ha de ser declarado ganador y cuál es la ordenación de los cinco candidatos según las preferencias del grupo en su conjunto? Los electores podrían intentar salir del *impasse* votando el método a emplear, pero ¿qué sistema de votación emplearían para decidirlo? No es inverosímil que reapareciera el mismo problema a este nivel superior, pues quizá los electores votaran por el método que más favoreciera a su candidato preferido.

(Esta tendencia natural a adaptar el enfoque de un problema a los propios intereses me recuerda el consejo del viejo abogado a su defendido: «Cuando la ley esté de su parte, aporree con la ley. Cuando los hechos estén de su parte, aporree con los hechos. Y cuando ni la ley ni los hechos estén de su parte, aporree la mesa». Debería señalar también que el problema de decidir

quién vota es aún más espinoso que el de decidir el sistema de votación. En general, la gente quiere que la ley dé derecho al sufragio al máximo número posible de partidarios y que se lo niegue —o que al menos los desanime— al máximo número razonablemente posible de adversarios. Ejemplos de esto último son la oposición al sufragio femenino y el *apartheid*, mientras que la vieja costumbre de hinchar la urna electoral ilustra el primer caso. Y no se limita a sucios fraudes en elecciones municipales, sino que con distintas variantes puede tentar incluso a las personas más altruistas, independientemente de su orientación política. Los antiabortistas contabilizan los «votos» de los no nacidos y, a menudo, los ecologistas van más allá y apelan al apoyo «electoral» de generaciones futuras no concebidas todavía).

Por lo que respecta a los sistemas de votación, la situación no es siempre tan confusa como sugiere el ejemplo anterior. Los números del ejemplo (debido a William F. Lucas por vía de los filósofos del siglo XVIII Jean-Charles de Borda y el marqués de Condorcet, así como de otros teóricos posteriores) fueron preparados para demostrar que el método de votación empleado puede determinar a veces el ganador. Pero aunque esas anomalías no se presenten siempre, cualquier método de votación está sujeto a ellas.

De hecho, el economista matemático Kenneth J. Arrow ha demostrado que no hay un procedimiento infalible para determinar las preferencias de un grupo a partir de las preferencias individuales que garantice el cumplimiento simultáneo de estas cuatro condiciones mínimas: si el grupo prefiere X a Y e Y a Z, entonces prefiere X a Z; las preferencias (tanto individuales como colectivas) han de limitarse a las alternativas disponibles; si todos los individuos prefieren X a Y entonces también el grupo prefiere X a Y; y no hay ningún individuo cuyas preferencias determinen dictatorialmente las del grupo.

Aunque todo sistema de votación tiene consecuencias indeseables y aspectos defectuosos, algunos sistemas son mejores que otros. Uno que quizá sería especialmente apropiado para unas primarias presidenciales, en las que se presentan varios candidatos, se llama votación de aprobación. En este sistema, cada elector puede votar por, o aprobar, tantos candidatos como quiera. El principio de «una persona, un voto» se sustituye por el de «un candidato, un voto» y se proclama vencedor el candidato que recibe la

máxima aprobación. No se producirían así situaciones como la de dos candidatos liberales que dividen el voto liberal y permiten que gane un candidato conservador con sólo el 40% de los votos.

El mandato moral de ser demócrata es formal y esquemático. La cuestión de fondo es cómo deberíamos ser demócratas y el enfocar esta cuestión con una actitud experimental abierta es perfectamente compatible con un firme compromiso con la democracia. A los políticos que, beneficiándose de un sistema electoral particular y limitado, se envuelven con el manto de la democracia, hay que recordarles de vez en cuando que este manto se puede presentar en varios estilos, todos ellos con remiendos.

Sustituibilidad y más sobre rutina

Conozco personas inteligentes capaces de seguir con facilidad los argumentos legales más complicados, las discusiones sentimentales más cargadas de matices y los relatos históricos más enrevesados, pero que, cuando se enfrentan a un sencillo «problema de letra» de matemáticas, se les pone inmediatamente una mirada vidriosa de no entender nada. Se quedan helados y se olvidan de plantearse las preguntas heurísticas de sentido común adecuadas, como suelen hacer en otros dominios de la vida: ¿de dónde viene este problema? ¿Qué quiero encontrar realmente y por qué? ¿Cómo podría simplificar la situación u obtener una respuesta aproximada? ¿Tiene el problema algo que ver con alguna cosa que conozca ya? ¿Puedo proceder en sentido inverso, de la solución a los datos?

A muchas de estas personas les parece que para los problemas matemáticos hace falta un modo de pensar rutinario y estúpido, así como la realización instantánea de algún tipo de cálculo. Si no se dan inmediatamente de narices con la respuesta, les parece que nunca van a hallarla. Encontrarían raro abordar un discurso matemático, pensar un problema matemático con palabras (véase también la entrada sobre *Cálculo rutina*). Como el personaje de Molière, que se sorprendió al descubrir que había estado toda la vida hablando en prosa, muchas personas se sorprenden cuando les dicen que buena parte de lo que llaman sentido común o lógica no es otra cosa que matemática.

Estas actitudes erróneas quizá procedan en parte de que el discurso de la matemática formal tiene algunas propiedades peculiares que no tienen los argumentos legales, las conversaciones sentimentales o los relatos históricos. Para ilustrar un ejemplo, que no por pequeño es menos importante, volvamos

a Euclides y a la venerable práctica de sustituir iguales por iguales.

Si nos acomete el impulso perverso de hacerlo, en un cálculo siempre podemos sustituir «25» por « 5^2 » o por « $(3^3 - 2)$ » sin que ello afecte al resultado. Análogamente, si en una discusión matemática sustituimos en todas partes «triángulo equilátero» por «triángulo equiángulo», la discusión seguirá teniendo el mismo sentido, pues las dos expresiones son maneras distintas de referirse a la misma clase de figuras. En general, las expresiones que adoptemos para describir o denotar los objetos matemáticos no afectan a la veracidad de los enunciados en los que aparecen estas expresiones. Esta propiedad de sustituibilidad, que se conoce habitualmente como extensividad, puede parecer perfectamente razonable y obvia, pero es característica de la matemática formal.

Decimos que dos conjuntos de números son iguales si tienen los mismos elementos, pues el modo en que se describen los conjuntos matemáticos no tiene importancia. Por contra, decimos que dos clubs escolares son distintos aunque tengan precisamente los mismos estudiantes, pues es esencial la caracterización (los objetivos) del club. En lógica se dice que la conversación informal sobre creencias, aspiraciones, objetivos y otras cosas por el estilo es comprensiva y no extensiva, y que por esta razón no permite la sustituibilidad. Por ejemplo, si alguien de la costa este con unas ideas geográficas no demasiado claras *cre*e que Cheyenne está en Montana, aunque «Cheyenne» sea igual a «la capital del estado de Wyoming», ciertamente no se sigue de ahí que este alguien crea que la capital del estado de Wyoming está en Montana. La primera caracterización de la ciudad no se puede sustituir por la segunda en este contexto comprensivo de creencia.

El programa de ordenador que traduce «El espíritu está dispuesto, pero la carne es débil» por «El vodka es agradable, pero la carne es tierna», o «tres hombres sabios» por «tres tipos listos» está atribuyendo a los lenguajes naturales (y en particular a los proverbios) una extensividad de la que simplemente carecen. No podemos decir que una familia que planea llegar a Disneylandia el 7 de enero esté planeando llegar para el cumpleaños de Millard Fillmore, aunque resulte que «7 de enero» y «cumpleaños de Millard Fillmore» denoten el mismo día. Como antes, la sustitución de iguales por iguales no conserva la veracidad del enunciado.

Quiero recalcar, no obstante, que la distinción entre contextos extensivos y comprensivos no es equivalente a la distinción entre contextos matemáticos y no matemáticos. He puesto mucho cuidado antes en escribir que sólo la matemática formal es extensiva; cuando se manejan símbolos, se comprueban deducciones o se hacen cálculos, no hay nada que objetar a la sustitución. Pero es seguro que en la interpretación y la aplicación de la matemática se habla también de necesidades, creencias y objetivos y, en estos contextos más humanos, la matemática es comprensiva también. Buena parte del estudio matemático, ya sea en un ambiente profesional o en la vida cotidiana, se dedica a aprender cómo demostrar los teoremas de la matemática formal, cómo interpretarlos en una situación concreta, y cómo y cuándo aplicar las reglas y fórmulas obtenidas. Aquí hay valores e historias —el origen del problema, su relación con otros problemas, sus posibles aplicaciones— y aquí es donde aparecen la heurística, la charla informal y los contextos comprensivos. En estos aspectos, el pensamiento matemático se parece mucho más al derecho, la historia, la literatura y la vida diaria.

Cada progreso matemático importante tiene su historia que le da contenido y significado: el teorema de Pitágoras, el desarrollo de nuestro sistema numérico, los avances de los árabes en álgebra, la evolución del cálculo desde Isaac Newton hasta Leonhard Euler, la geometría no euclídea, la teoría de Galois, el teorema de Cauchy en análisis complejo, la teoría de conjuntos de Cantor, el teorema de incompletitud de Gödel, y muchísimos otros teoremas e ideas. ¿No son más que simples cálculos o demostraciones formales? Normalmente se cree que los cálculos y las demostraciones son lo característico de la matemática, pero por necesarios que sean a veces, la mayoría de gente tampoco los quiere la mayor parte del tiempo. Quieren lo mismo que en otros campos: explicaciones, historias y heurística.

La idea de que el modo de pensar matemático difiere en lo fundamental del de otros campos es perniciosa. Uno de sus orígenes reside en los profesores de matemáticas que no saben establecer la conexión entre lo que enseñan y el resto del plan de estudios, ni con sucesos y noticias corrientes que permitan un enfoque matemático, ni con ningún otro aspecto de la vida cotidiana del estudiante. Otra razón es la imagen equivocada de la frialdad, la irrelevancia y la dificultad de las matemáticas. Y otra razón más, como ya he

sugerido aquí, es el malentendido filosófico que confunde la sustituibilidad y la extensividad vacía de la matemática formal con el núcleo central de esta materia, más rico y comprensivo (y además realza a aquélla frente a éste). La matemática tiene tanto de narración, propósitos y relatos como de cálculo y fórmulas. Si no somos capaces de darnos cuenta de ello y permanecemos en la ignorancia de la matemática pero ciegamente reverentes hacia sus técnicas, nos empobrecemos sin necesidad y delegamos demasiado en otros.

Tautologías y tablas de verdad

O Aristóteles era pelirrojo o Aristóteles no era pelirrojo. Como no es verdad que uno de los dos, Gottlob o Willard, esté presente, entonces los dos, Gottlob y Willard, están ausentes. Siempre que Thoralf está fuera de la ciudad Leopold vomita, por tanto si Leopold no está vomitando, entonces Thoralf no está fuera de la ciudad. Todas estas genialidades matemáticas son ejemplos de tautologías, enunciados que son verdaderos en virtud del significado de los conectivos lógicos «no», «o», «y», y «si..., entonces...». («Tautología» se usa también informalmente en un sentido más amplio).

Si simbolizamos por letras mayúsculas las frases afirmativas simples que componen estos enunciados, podemos expresarlos como: «A o no A», «Si no (G o W), entonces (no G y no W)» y «Si (si T, entonces L), entonces (si no L, entonces no T)». Podemos sustituir A, G, W, T y L por cualquier otra declaración simple y el resultado seguirá siendo cierto. Así, «O Gorbachov es un travestido o no lo es», «Como no es cierto que uno de los dos, Jorge o Marta, sea culpable, entonces los dos, Jorge y Marta, son inocentes», y «Como siempre que llueve la tienda de ordenadores está cerrada, si la tienda de ordenadores está abierta, entonces no llueve» son enunciados tautológicamente verdaderos los tres. Estas tautologías en particular tienen además un nombre y se llaman ley del medio excluso, ley de Morgan y ley de contraposición, respectivamente.

Los lógicos han formalizado el proceso de comprobación por el que se determina el carácter tautológico de estas proposiciones. Han inventado reglas, que generalmente se llaman tablas de verdad, para cada uno de los conectivos: una proposición «No P» es verdadera precisamente cuando P es falsa; las proposiciones «P y Q» son verdaderas sólo si P y Q son verdaderas

a la vez; las proposiciones « P o Q » son verdaderas sólo cuando al menos una de las dos lo es; y las proposiciones «Si P , entonces Q » sólo son falsas cuando P es verdadera y Q es falsa. (Habría que decir que hay otros usos no matemáticos del «si..., entonces...» que se interpretan de un modo distinto. Pero desde el punto de vista matemático la proposición «Si la luna estuviera hecha de queso verde, entonces Bertrand Russell sería Papa» es verdadera y, en general, es verdadero afirmar que de un enunciado falso se sigue cualquier cosa). Y, finalmente, los enunciados « P si y sólo si Q » son verdaderos únicamente cuando P y Q tienen el mismo valor de verdad, o son las dos verdaderas o las dos falsas.

Dos pequeñas digresiones. ¿Qué dice en realidad el siguiente anuncio de una plaza de profesor de matemáticas? «Buscamos un candidato que sea un profesor entusiasta y eficiente o que haga investigación. Desgraciadamente no podemos considerar aquellos candidatos que, siendo profesores entusiastas y eficientes, no hagan también investigación». Dejando ahora la vida académica, imagínese en una isla habitada por gentes que o siempre dicen la verdad o siempre mienten. Usted se encuentra ante dos carreteras y quiere saber cuál de las dos lleva a la capital. Por suerte un habitante del lugar pasa por allí (usted no sabe si es de los mentirosos o de los que dicen la verdad) y sólo tiene tiempo de contestarle sí o no a una pregunta. ¿Qué pregunta ha de hacerle usted para determinar cuál es la carretera que lleva a la capital?

En lógica proposicional, que es como se llama la parte de la lógica matemática que estoy describiendo, lo contrario de una tautología es una contradicción, la frase que es falsa en virtud del significado de sus conectivos lógicos. Podemos decir que «Juan es calvo y Juan no es calvo» es falsa sin saber nada de Juan ni de su pelo. Las contradicciones formales tales como « A y no A » y «(C y no B) y (si C , entonces B)» son falsas independientemente de las frases simples que pongamos en lugar de A , B y C . La tercera categoría de enunciados en lógica proposicional, la mayor de todas, comprende aquellos que a veces son verdaderos y a veces son falsos según sea la verdad o falsedad de sus constituyentes. Sólo como ejemplo de tabla de verdad, en la página siguiente consideraremos el caso de « A y (B o no A)».

En las columnas de A y B tenemos los cuatro argumentos de verdad y falsedad posibles para este par de símbolos. La primera fila se ha de

interpretar como que el enunciado «A y (B o no A)» es globalmente verdadero (lo indica la V subrayada) siempre que A y B sean ambos verdaderos. Las tres filas restantes dicen que el enunciado es falso para cualesquiera otros valores de verdad asignados a A y B. [Los símbolos tradicionales de «y», «o» y «no» son $\wedge$, $\vee$, y $\neg$, con lo que la proposición anterior se expresa $A \wedge (B \vee \neg A)$. Una flecha $\rightarrow$, es el símbolo de «si..., entonces...», mientras que una flecha de dos puntas $\leftrightarrow$, indica «si y sólo si»].

A	B	A	y	(B	o	no	A)
V	V	V	<u>V</u>	V	V	F	V
V	F	V	F	F	F	F	V
F	V	F	F	V	V	V	F
F	F	F	F	F	V	V	F

En realidad, para determinar la verdad o falsedad de esta proposición (que resulta tener los mismos valores de verdad que «A y B») no hacen falta tablas de verdad, pero éstas son a menudo valiosísimas para afirmaciones más complicadas que contienen paréntesis anidados (fórmulas dentro de fórmulas) y más símbolos de frase que los dos de este ejemplo. Aparte de determinar si una proposición dada es una tautología, una contradicción o una afirmación contingente, las tablas de verdad se pueden usar para determinar la validez de ciertos tipos de razonamiento hechos famosos por Lewis Carroll. (Un ejemplo sacado de la economía del país de las maravillas: si sube el mercado de obligaciones o bajan los tipos de interés, entonces cae la bolsa o no suben los impuestos. La bolsa cae si y sólo si sube el mercado de obligaciones y aumentan los impuestos. Si bajan los tipos de interés, entonces la bolsa no cae o el mercado de obligaciones no sube. Por tanto, o aumentan los impuestos o cae la bolsa y bajan los tipos de interés). Los protocolos para comprobar esta clase de validez son tan rutinarias que todos los ordenadores llevan incorporados a su *hardware* circuitos para «y», «o» y «no», de modo que las máquinas puedan detectar casi instantáneamente la verdad de sentencias y condiciones complejas.

Pero los mecanismos de tabla de verdad sirven de bien poco para proposiciones que contienen frases relacionales. (Véase la entrada sobre *Los cuantificadores*). La afirmación «Todos los amigos de Mortimer son amigos

míos», «Oscar es amigo de Mortimer» y «Oscar es mi amigo», consideradas como enunciados de lógica proposicional, han de simbolizarse con letras — pongamos P, Q y R, respectivamente—. Estos símbolos no reflejan el hecho de que P y Q impliquen R pues la implicación no depende de los significados de «y», «o», «no» y «si..., entonces...». La implicación sólo es captada en la lógica predicativa, que además de la lógica proposicional abarca la lógica de las frases relacionales («es amigo de» en este caso) y los cuantificadores asociados (aquí el «todo»). En este dominio más amplio el lógico norteamericano Alonzo Church ha demostrado que nunca puede haber una receta (como el método de las tablas de verdad) para determinar la validez de las frases o razonamientos.

[Soluciones a los acertijos: si se rompe un poco la cabeza se convencerá de que las condiciones para el puesto se reducen a hacer investigación. El cálculo da el mismo resultado, pues cualquiera que sea el valor que asignemos a E (profesor entusiasta y eficiente) y a I (hacer investigación) $(E \vee R) \wedge \neg(E \wedge \neg R)$ tiene el mismo valor de verdad que R. Y una buena pregunta a plantear al habitante que dice la verdad o es mentiroso es «¿Es verdad que la carretera de la izquierda lleva a la capital si y sólo si usted dice siempre la verdad?». Si la carretera de la izquierda lleva a la capital, tanto los que dicen la verdad como los mentirosos contestarán «sí», y «no» en caso contrario. Otra pregunta podría ser: «Si le preguntara si la carretera de la izquierda lleva a la capital ¿me diría usted que sí?». Aquí también, los mentirosos y los que dicen siempre la verdad darían la misma respuesta].

El teorema de Pitágoras

Aunque sea discutible, suele decirse que los primeros matemáticos fueron Pitágoras (hacia 540 a. C.) y su antecesor Tales (hacia 585 a. C.). Existieron, por supuesto, pueblos anteriores que poseían unos notables conocimientos matemáticos (el papiro de Rhind del 1650 a. C. es un filón impresionante de instrumentos de cálculo, entre los que se incluye el hábil uso de una notación rudimentaria para las fracciones), pero los egipcios, babilonios y demás pueblos tenían una actitud muy distinta hacia esta materia. Para ellos, la matemática sólo era una materia útil para fijar impuestos, calcular intereses, determinar el número de cuarteras de cebada que se necesitaban para hacer una cierta cantidad de cerveza, calcular las áreas de los campos, los volúmenes de sólidos, las cantidades de ladrillos y los datos astronómicos. Aunque no cabe la menor duda de que se trata de técnicas importantes, que tristemente superan la capacidad de demasiados ciudadanos contemporáneos, desde el tiempo de Pitágoras y Tales la matemática ha significado algo más que mero cálculo.

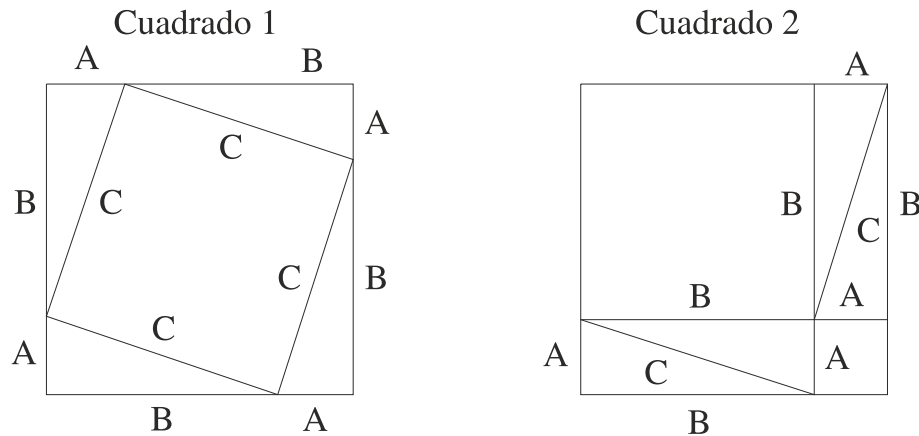
Nunca antes del siglo VI a. C. había considerado nadie la matemática como algo que tuviese una estructura lógica, como algo que admitiera una sistematización racional, o como un conjunto de conceptos ideales que podían aclararse mediante la aplicación de la razón humana. Tales, Pitágoras y sus contemporáneos y discípulos se dieron cuenta de ello. Nadie antes que ellos vio los números y las formas geométricas como algo omnipresente, ni tampoco nadie pensó en términos de círculos teóricos y números abstractos en lugar de ruedas de carro y números concretos. Tales y Pitágoras, sí. Nadie pensó en extraer las realidades más elementales y evidentes relativas a estos conceptos matemáticos y luego, a partir de estas realidades fundamentales,

intentar obtener otras, teoremas menos evidentes, mediante la sola lógica. Coro: Tales y Pitágoras, sí. Ellos y los matemáticos griegos que les siguieron inventaron la matemática (y la lógica) tal y como la conocemos; la fundaron como arte liberal y no como un simple mascar números.

De Pitágoras nos han llegado pocos datos personales. Viajó mucho, fundó la sociedad mística de los pitagóricos que prohibía comer habas y cuyo lema era «Todo es número», y algunos dicen que acuñó las palabras «filosofía» («amor a la sabiduría») y «matemático» («estudioso»). Pitágoras y sus discípulos tuvieron una gran influencia sobre la matemática griega (es decir, sobre la matemática) y se les atribuye haber descubierto gran parte de lo que 250 años después constituiría el primero de los dos libros de los *Elementos* de Euclides, y en particular el teorema que invariablemente lleva su nombre. (Véase la entrada sobre *Geometría no euclídea*). El teorema de Pitágoras es uno de los importantes e indispensables, y sus demostraciones más corrientes han sido ejemplos de belleza geométrica durante casi veinticinco siglos. Dice que en un triángulo rectángulo, el área del cuadrado construido sobre la hipotenusa (o lado opuesto al ángulo recto) es igual a la suma de las áreas de los cuadrados construidos sobre los otros dos lados. O escribiéndolo en una forma más simbólica, si las longitudes de estos dos lados son A y B y la longitud de la hipotenusa es C, entonces podemos asegurar que $C^2 = A^2 + B^2$. [He aquí una pista délfica para una demostración del teorema: hay dos maneras distintas de colocar cuatro triángulos rectángulos de lados A, B y C en el interior de un cuadrado de lado (A + B). Una de ellas deja al descubierto una superficie cuadrada de lado C, y la otra deja sin cubrir dos porciones cuadradas de lados A y B, respectivamente].

Aunque Pitágoras probablemente habría tenido poco interés en ello, su teorema puede emplearse para calcular distancias. Así, si Elvira está 12 kilómetros al norte del Partenón y Guillermo está 5 kilómetros al este de dicho edificio, podemos calcular que, en línea recta, Elvira está exactamente a 13 kilómetros de Guillermo, pues $5^2 + 12^2 = 13^2$. De modo análogo se puede determinar la longitud de la diagonal de un rectángulo o de una caja de zapatos. Expresado en el lenguaje de la geometría analítica (que no sería descubierta hasta 2000 años después) y convenientemente generalizado, el

teorema de Pitágoras es un instrumento matemático potentísimo.



Las áreas de los cuadrados 1 y 2 son iguales. Las áreas de los cuatro triángulos del cuadrado 1 y las de los mismos triángulos del cuadrado 2 son iguales. Por tanto, el área restante interior al cuadrado 1 es igual al área restante interior al cuadrado 2. Esto es, $C^2 = A^2 + B^2$

No obstante, Pitágoras habría comprendido mejor la reacción estética del filósofo Thomas Hobbes ante su teorema. Un amigo de Hobbes, John Aubrey, escribió que el filósofo tenía cuarenta años cuando por casualidad echó una ojeada a un libro de geometría. Estaba en la biblioteca de un caballero y se encontró con los *Elementos* de Euclides abierto sobre una mesa, y «era el teorema de Pitágoras. Hobbes leyó la proposición. “Por D___”, exclamó. (De vez en cuando juraría, a modo de énfasis). “Por D___”, exclamó, “¡esto es imposible!”. Y leyó la demostración, que hacía referencia a tal otra proposición anterior, que a su vez leyó también. Ésta le dirigió a otra anterior, y también la leyó. *Et sic deinceps*, que por fin se convenció demostrativamente de la verdad. Esto hizo que se enamorara de la geometría».

He de confesar que la primera vez que estudié geometría me entró una chifladura de la misma clase (aunque resistí a la tentación de ir repitiendo «Por D___»). Desgraciadamente, ya no se comunica a los estudiantes este inestimable legado que es el método axiomático, la deducción de proposiciones no intuitivas a partir de axiomas evidentes. Demasiados textos de geometría parecen haber optado por un enfoque de la disciplina anterior al

de los griegos, y prácticamente sólo ponen el acento en hechos desconectados, reglas empíricas y fórmulas prácticas. Pitágoras habría preferido comer habas a leer alguno de estos textos.

Teoría de juegos

Prácticamente todas las situaciones de la vida se pueden considerar como juegos si interpretamos la palabra «juego» con la suficiente laxitud. (Naturalmente, si interpretamos «suficiente laxitud» con suficiente laxitud, muchas situaciones de la vida se pueden considerar también como melones, pero esto sería un exceso que superaría la tolerancia lingüística de cualquiera). No es pues sorprendente que la teoría matemática de juegos tenga un papel esencial en el método empleado por los planificadores económicos, militares y políticos para enmarcar sus elecciones y decidir sus estrategias. Su inventor fue John von Neumann hace unos cincuenta años y tenía en mente estas aplicaciones, y puede servir también para aclarar el significado de ciertas decisiones personales y de algunos compromisos.

La teoría de juegos es muy útil cuando interviene el farol o *bluff* y es, por tanto, necesario recurrir a estrategias probabilísticas. En juegos con una información perfecta, como las damas o el ajedrez, siempre hay una estrategia determinista óptima, y los movimientos no tienen por qué ser aleatorios ni secretos. Aunque se sabe mucho acerca de los juegos de esta clase, el hecho de que exista una estrategia vencedora no implica que ésta pueda encontrarse en «tiempo real». Hasta el momento no se conocen las estrategias óptimas del ajedrez ni de las damas, pero las de juegos más simples, como el de las tres en raya, son bien conocidas por las educadoras de guardería.

Una situación de juego se da cuando hay dos o más jugadores, y cada uno de ellos es libre de escoger entre un conjunto de posibles opciones o estrategias. Cada elección comporta a su vez un resultado distinto —premios o sanciones de mayor o menor importancia—. Cada jugador tiene sus

preferencias con respecto a tales resultados. La teoría de juegos se ocupa de determinar las estrategias de los jugadores, sus costes y ganancias, y las situaciones de equilibrio.

Pero en vez de exponer los principios de la materia, describiré una situación de juego típica que se presta a estrategias probabilísticas. Consideremos un lanzador (*pitcher*) frente a un bateador en un partido de béisbol. El lanzador puede mandar una pelota curvada (*curveball*), una pelota rápida (*fastball*) o una pelota con efecto (*screwball*). Las medias de acierto del bateador son las siguientes. Si se espera una pelota rápida, su media es de 0,300 si recibe una curvada (esto es, acierta el 30% de las veces), 0,400 si recibe una pelota rápida y 0,200 si recibe una con efecto. Pero si se espera una pelota curvada, sus medias son 0,400 para estos lanzamientos, 0,200 para las rápidas y 0,000 para pelotas con efecto. Y si lo que se espera es una pelota con efecto, sus medias para las curvadas, rápidas y con efecto son, respectivamente, 0,000, 0,300 y 0,400.

		El bateador espera		
		Pelota curvada	Pelota rápida	Pelota con efecto
El lanzador manda	Pelota curvada	0,4	0,3	0
	Pelota rápida	0,2	0,4	0,3
	Pelota con efecto	0	0,2	0,4

Probabilidades de acierto

En base a estas probabilidades, el lanzador ha de decidir qué tipo de lanzamiento hará y el bateador, previéndolo, ha de prepararse para recibirlo. Si el bateador se prepara para una pelota rápida evitará efectivamente tener una media de acierto del 0,000. Pero si hace esto constantemente, el lanzador sólo le mandará pelotas con efecto, y le mantendrá en su media de aciertos más baja de 0,200. El bateador podría entonces decidir estar preparado para las pelotas con efecto, que, si el lanzador insiste en ellas, le darían una media de aciertos de 0,400. El lanzador podría entonces prever esto y lanzar pelotas curvadas que, si el bateador sigue esperando con efecto, dejarían la media de aciertos en un 0,000. Está claro que este modo de razonar podría seguir cíclica e indefinidamente.

Cada jugador tiene que idear una estrategia probabilística general. El

lanzador ha de decidir qué porcentaje de sus lanzamientos han de ser pelotas curvadas, rápidas y con efecto y luego lanzarlos *al azar* según estos porcentajes. Y, por su parte, el bateador ha de decidir también qué porcentaje de veces ha de esperar cada tipo de lanzamiento y luego estar preparado *al azar* de acuerdo a estos porcentajes. Las técnicas y teoremas de la teoría de juegos nos permiten encontrar las estrategias óptimas para cada jugador en este juego así como en otros muchos. La solución para este juego ideal en particular resulta ser: para el lanzador enviar un 60% de pelotas con efecto y un 40% de curvadas, y para el bateador estar preparado para recibir pelotas rápidas un 80% de las veces y con efecto el 20% restante. Si ambos siguen estas estrategias óptimas, la media de aciertos del bateador será de 0,240.

El conocido «juego del gallina» no admite una solución tan clara. En una de sus versiones intervienen dos adolescentes que dirigen sus coches uno contra el otro a gran velocidad. El primero que se desvía pierde prestigio y el otro es el vencedor. Si ambos se desvían, empatan. Y si no lo hace ninguno de los dos, chocan. En términos más cuantitativos, los adolescentes A y B tienen la posibilidad de elegir entre girar el volante o no hacerlo. Si A gira el volante y B no, supondremos a título de ilustración que el resultado es 20 puntos para A y 40 para B. Si B gira y A no, la puntuación es la inversa. Si ambos se desvían, las ganancias son 30 puntos cada uno, mientras que si ninguno de los dos se desvía la «ganancia» es de 10 puntos cada uno. Como en el dilema del preso (véase la entrada sobre *Ética y matemáticas*), se trata de una situación bastante general y no es exclusiva de los adolescentes cretinos. Y, como en el dilema del preso, también se da el hecho de que si un individuo persigue sólo maximizar sus beneficios personales, no lo consigue.

No hace falta un gran derroche de imaginación para darse cuenta de que hay muchas situaciones en los negocios (conflictos laborales y guerras por mercados), el deporte (prácticamente todos los concursos competitivos) y el ámbito militar (juegos de guerra) que pueden modelizarse en términos de la teoría de juegos. Un ejemplo relativamente nuevo lo tenemos con los aparatos que revelan el número de teléfono de la persona que le está llamando, la opción del que llama de impedir que usted descubra su número y la opción de usted en cada caso de contestar o no a la llamada. Aunque la mayoría de aplicaciones suelen emplear palabras inquietantes como «batalla», «guerra» y

«contienda», este vocabulario no es esencial. El asunto igual podría llamarse teoría de la negociación que teoría de juegos. Sus principios son aplicables en los llamados juegos de suma no nula en los que las ganancias de un jugador no tienen por qué ser equilibradas por las pérdidas de otro, en las negociaciones íntimas (¿guerra de los sexos?), en juegos más «comunales» como mantener un gobierno justo, e incluso en aquellos en los que uno de los jugadores es la naturaleza o el medio ambiente.

Otra cosa a tener en cuenta es que hay que evitar que el aparato técnico de la teoría de juegos nos haga perder de vista las suposiciones que intervienen implícitamente en una negociación o en una contienda concreta. («Tuvimos que destruir la ciudad para salvarla»). Desgraciadamente es muy fácil quedarse fascinado en la construcción de matrices de ganancias y en el cálculo de las consecuencias esperadas de varias estrategias, olvidando considerar las distintas suposiciones que nos permiten hacerlo, así como de los objetivos que se persiguen. Personalmente padezco de estos efectos anestésicos de la tecnofilia y este libro lo escribí en parte como expiación.

La teoría del caos

La gente nos acusa a menudo de que el conocimiento de la matemática produce una ilusión de certeza y la consiguiente arrogancia. Creo que esto es falso y, a modo de ejemplo en el que ocurre exactamente lo contrario, recurro a la teoría del caos. El nombre de este campo del conocimiento, como el de su prima, la teoría de catástrofes, parece especialmente adecuado para una creación matemática del siglo XX, pero olvidémoslo por unos instantes. Todos los dominios técnicos suelen apropiarse de palabras corrientes y las retuercen para convertirlas en parodias deformadas de sí mismas.

La teoría del caos no versa sobre tratados anarquistas ni sobre manifiestos surrealistas o dadaístas, sino sobre el comportamiento de sistemas no lineales arbitrarios. Para lo que nos interesa, un sistema se puede considerar como un conjunto de partes cuyas interacciones se rigen por unas determinadas reglas y/o ecuaciones. El Servicio Postal, el sistema circulatorio humano, la ecología local y el sistema operativo del ordenador que estoy usando, son ejemplos de este vago concepto de sistema. Hablando también en un tono completamente impreciso, un sistema no lineal es aquel cuyas partes no están relacionadas de un modo lineal o proporcional (como lo están, por ejemplo, una báscula de baño o un termómetro); doblar la magnitud de una parte no significa que se vaya a doblar la de otra, y la salida tampoco es proporcional a la entrada. Para entender mejor estos sistemas, la gente se vale de modelos —maquetas a escala reducida, formulaciones matemáticas y simulaciones por ordenador— con la esperanza de que estos modelos más simples arrojen alguna luz sobre los sistemas en cuestión.

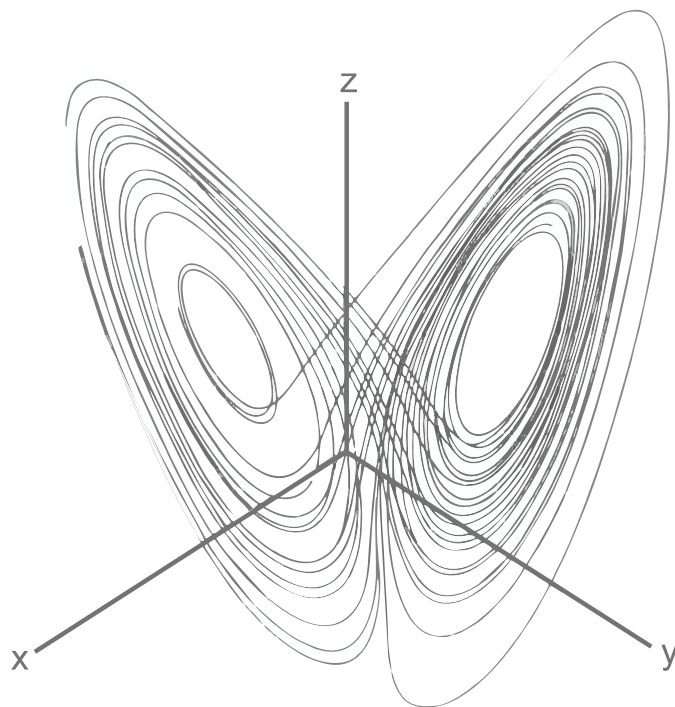
En 1960, jugando con uno de esos modelos por ordenador de un sistema meteorológico sencillo, Edward Lorenz descubrió algo muy extraño.

Introduciendo inadvertidamente en su modelo unos datos que diferían entre sí en menos de una milésima, descubrió que las proyecciones meteorológicas resultantes divergían enseguida cada vez más, hasta que no guardaban ninguna relación entre sí. Éste fue, según ha observado el autor James Gleick, el inicio de la ciencia matemática de la teoría del caos.

Aunque el modelo no lineal de Lorenz era una simplificación y su equipo informático era muy primitivo, sacó las conclusiones correctas de la divergencia de los pronósticos simulados por ordenador: la causa eran las pequeñas variaciones en las condiciones iniciales del sistema. Más generalmente, los sistemas cuya evolución se rige por reglas y ecuaciones no lineales pueden ser sumamente sensibles a cambios tan ínfimos, y a menudo presentan un comportamiento impredecible y «caótico». Por el contrario, los sistemas lineales son mucho más robustos, y pequeñas diferencias en las condiciones iniciales sólo producen pequeñas diferencias en los resultados finales.

El tiempo, incluso en este modelo simplificado, no admite pronósticos a largo plazo porque es demasiado sensible a cambios casi imperceptibles en las condiciones iniciales, los cuales producen cambios ligeramente mayores un minuto después o un metro más allá, y éstos a su vez dan lugar a desviaciones más notables, de modo que todo el proceso entra en una cascada de impredecibilidad no repetitiva. Si se sigue la evolución de este tiempo simulado, se observa que, si bien se cumplen ciertas condiciones generales (no hay ventiscas en Kenia, gradientes de temperatura estacionales abruptos, etc.), los pronósticos concretos a largo plazo, en el supuesto de que alguien tuviera la temeridad de hacerlos por un año o dos, carecen virtualmente de valor.

Desde el trabajo de Lorenz se han dado muchas manifestaciones del efecto mariposa, o sensibilidad de los sistemas no lineales a las condiciones iniciales, en disciplinas científicas que van desde la hidrodinámica (turbulencia y movimiento de un fluido) a la física (osciladores no lineales), o de la biología (fibrilaciones cardíacas y epilepsia) a la economía (fluctuaciones de los precios). Estos sistemas no lineales muestran una complejidad sorprendente que parece darse aun en el caso de que las reglas y ecuaciones que definen el sistema sean elementales.



Una trayectoria en un espacio de fases abstracto. La figura en forma de mariposa se conoce como atractor de Lorenz

Sin entrar demasiado en detalles, insistiré en que el estado de estos sistemas en un instante dado puede representarse por un punto en un espacio matemático de varias dimensiones que se llama «espacio de fases», y que su evolución subsiguiente se representa por una curva en este espacio de fases abstracto. Las trayectorias de estos sistemas resultan ser aperiódicas e impredecibles, la gráfica de todas las trayectorias posibles a menudo presenta un sinfín de circunvoluciones y si se la examina de cerca presenta aún más complejidad. Un examen todavía más preciso de las trayectorias del sistema en el espacio de fases revela la presencia de vórtices todavía más pequeños y de complicaciones del mismo tipo genérico. En pocas palabras, el conjunto de todas las trayectorias posibles de estos sistemas no lineales forma un fractal (véase la entrada sobre *Fractales*) y es una de las muchas factorías donde se producen esos monstruos geométricos del matemático Benoît Mandelbrot, que de pronto se han hecho omnipresentes.

Como ya advertí al principio, el descubrimiento de la teoría del caos podría fomentar en nosotros una cierta cautela. Una causa de esta falta de

confianza es el hecho de que estos sistemas a menudo se comportan de un modo perfectamente normal y suave para un amplio dominio de condiciones iniciales y de repente se vuelen caóticos cuando un determinado parámetro del sistema alcanza un valor crítico. Imaginemos, por poner un ejemplo singularmente simple debido al físico Mitchell Feigenbaum, que la población de una determinada especie animal viene dada por la fórmula no lineal $X' = RX(1 - X)$, donde X' es la población en un determinado año, X es la población en el año anterior, y R es un parámetro que varía entre 0 y 4. Por simplicidad supondremos que X y X' son números comprendidos entre 0 y 1, de modo que la población verdadera es 1.000.000 de veces estos valores.

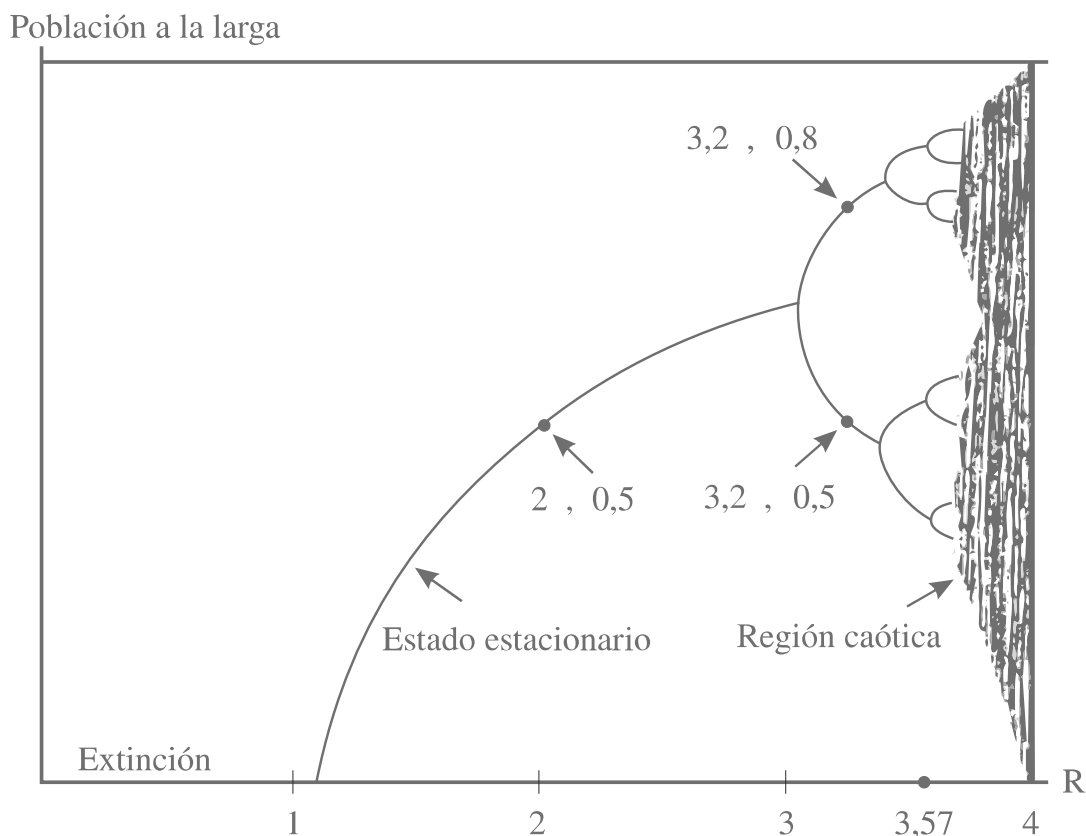
¿Cómo evoluciona la población de esta especie si $R = 2$? Si suponemos que la población X es ahora 0,3 (esto es, 300.000), aplicando la fórmula sabemos que la población será $2(0,3)(0,7) = 0,41$. Para obtener la población al año siguiente, introducimos 0,41 en la fórmula y encontramos $2(0,41)(0,59) = 0,4838$. Les ahorraré la aritmética, pero se puede emplear el mismo procedimiento para determinar la población al cabo de tres años, de cuatro, de cinco, de seis, etc. Encontramos que la población se estabiliza en 0,5. Es más, para cualquier valor inicial encontraríamos que la población también se estabiliza en 0,5. Esta población se llama población del estado estacionario para este valor de R .

Cálculos parecidos con un valor de R menor, pongamos $R = 1$, demuestran que sea cual sea el estado inicial, la población se «estabiliza» en 0, se extingue. Cuando R es mayor, pongamos 2,6 para concretar, encontramos después de aplicar este procedimiento que, independientemente de cual sea el estado inicial, la población se estabiliza ahora en aproximadamente 0,62, el estado estacionario para este R .

¿Y qué más? Bueno, pues aumentemos R otra vez, tomemos ahora un valor de 3,2. Como antes, no importa cuál sea la población inicial, cuando aplicamos iteradamente la fórmula $3,2X(1 - X)$, encontramos que la población de la especie no se estabiliza en un valor, sino que se instala finalmente en una alternancia de dos valores —aproximadamente 0,5 y 0,8—. Esto es, un año la población es 0,5 y el siguiente es 0,8. Aumentemos nuestro parámetro R hasta 3,5 y veamos qué sucede. La población inicial tampoco importa ahora, pero esta vez lo que hace la población a la larga es tomar

alternativamente cuatro valores distintos, aproximadamente 0,3, 0,83, 0,50 y 0,88 en años sucesivos. Si volvemos a aumentar ligeramente el valor de R , la población se instala en una alternancia regular de ocho valores distintos. Incrementos cada vez menores de R llevan a nuevos desdoblamientos del número de valores.

Luego, de repente, aproximadamente en $R = 3,57$, el número de valores se hace infinito y la población de la especie varía al azar de un año a otro. [Se trata, sin embargo, de una extraña clase de azar, pues resulta de iterar la fórmula $3,57X(1 - X)$, con lo que la sucesión de poblaciones está totalmente determinada por la población inicial.] Aún más extraño que esta variación caótica de la población de la especie es el hecho de que un ligero aumento de R vuelve a dar lugar a una alternancia regular de la población de un año para otro, y un nuevo incremento vuelve a producir una variación caótica. Estas alternancias ordenadas, seguidas de caos aleatorio, al cual suceden a su vez ventanas de comportamiento regular, dependen críticamente del parámetro R , que parece ser una medida grosera de la volatilidad del modelo.



Relación entre los valores de R , poblaciones y caos. La región caótica es muy intrincada. Este tipo de diagrama es debido al biólogo Robert May

Aunque un primer plano de la región caótica revela una complejidad inesperada, la variación anual de la población no da un fractal especialmente bello como ocurre con algunos sistemas no lineales. Sin embargo, hay suficiente complejidad para dejar clara la siguiente enseñanza. Si un sistema tan trivial como esta única ecuación no lineal puede dar lugar a una tal impredecibilidad caótica, quizá no debiéramos ser tan enérgicos ni tan dogmáticos al referirnos a los resultados predichos por las distintas políticas sociales, económicas y ecológicas aplicadas a enormes sistemas no lineales como Estados Unidos o el planeta Tierra. Las consecuencias de esas políticas son, podría pensarse, bastante más oscuras que las del valor de R en este modelo.

Siempre es peligroso, y a menudo estúpido, aplicar resultados técnicos fuera de su contexto originario. Así pues, no pretendo sugerir con mis observaciones cautelares que debamos adoptar una actitud pasiva de lavarnos las manos por el mero hecho de que nunca podamos estar seguros de los efectos de nuestras acciones. Sólo pretendo decir que la teoría del caos (y muchas otras cosas) aconseja que la adopción de cualquier medida política, económica o militar vaya acompañada de un cierto escepticismo y una cierta humildad. (Por falta de un clavo se perdió la batalla y cosas así). Puedo añadir una anécdota personal al respecto: en algunas entrevistas destinadas a promocionar mi libro *El hombre anumérico* me preguntaron muchas veces qué opiniones globales tenía yo acerca de la decadencia de Estados Unidos. He de admitir que sucumbí algunas veces y me aventuré más allá de la educación matemática para caer en la monótona letanía del estado lamentable de nuestra sociedad. Sin embargo, la mayoría de las veces pude resistir, y contestaba que si no fui capaz de predecir el éxito de *El hombre anumérico*, mucho menos iba a serlo de la caída de la civilización norteamericana.

Se podría extraer la consecuencia de que estos sistemas no lineales y caóticos son animales raros, curiosidades que sólo preocupan a los matemáticos y a algunos científicos fanáticos. Esto no es ni muchísimo menos cierto. Como dijo en cierta ocasión el matemático Stanislaw Ulam,

llamar a la teoría del caos «ciencia no lineal» es como llamar a la zoología «estudio de los no elefantes». Ahora que disponemos de más útiles, los científicos han empezado a darse cuenta de que ya no es necesario esconder el ruido, la estática y el rozamiento ni ignorar las turbulencias, la arritmia ni las contingencias «aleatorias». Estos fenómenos, al igual que los fractales, están por todas partes y, por el contrario, son sus primos lineales los que, como los elefantes, son raros.

Test de Turing y sistemas expertos

El matemático y lógico inglés Alan M. Turing fue el autor de una serie de artículos fundacionales sobre lógica e informática teórica. En un artículo de 1936, antes de que se hubiera construido ningún ordenador programable, describió simbólicamente la estructura lógica que habría de tener cualquiera de esas máquinas. Su descripción de un ordenador ideal especificaba en términos matemáticos las relaciones entre la entrada, la salida, las acciones y los estados de lo que se ha dado en llamar una máquina de Turing. En otro artículo daba las razones por las que era irrelevante que el substrato físico de dicha máquina fueran neuronas, chips de silicio o un mecanismo de hojalata. Lo esencial es la pauta y la estructura, y no el material que hay debajo.

Durante la segunda guerra mundial, Turing trabajó en criptografía y en la tarea de descifrar los códigos secretos alemanes. Después volvió al trabajo abstracto y en un famoso artículo de 1950 propuso que la pregunta etérea de si podía considerarse que los ordenadores tenían conciencia fuera sustituida por la menos metafísica de si se puede programar un ordenador de modo que «engañe» a una persona y le haga creer que está tratando con otro ser humano. Mediante un monitor de televisión alguien podría plantear a un ordenador programado a propósito y a otra persona preguntas que se contestaran con un sí o un no, o con respuestas seleccionadas de una propuesta múltiple. Se trataría entonces de que una persona adivinara qué conjunto de respuestas era dado por el humano y cuál por el ordenador. Si esa persona fuera incapaz de distinguirlos, el ordenador habrá superado el test de Turing. A menudo el test se plantea en términos de conversaciones. Imagínese sosteniendo una conversación con dos interlocutores por medio de

un monitor. Su tarea consistiría en decidir cuál de los dos tiene un *hardware* (o fisiología) basada en el silicio y cuál en el carbono. (Con respecto a esto último, el filósofo norteamericano Hilary Putnam ha desarrollado una sugerente analogía entre la distinción que se hace en informática entre el *software* y el *hardware*, y la distinción entre el cerebro y la mente en filosofía).

De todos modos, los criterios para superar el test de Turing son muchísimo más claros que los que se refieren a la conciencia de una máquina. Pero a pesar de que Turing predijera que hacia el año 2000 habría máquinas capaces de superar su test, ninguna se ha aproximado siquiera a ello. Y ciertamente en lo que respecta al futuro inmediato, una «conversación» de ordenador revelaría enseguida su alma de metal. La cantidad de conocimiento tácito que poseemos todos nosotros supera con creces la capacidad de nuestros presuntos imitadores. Sabemos que los gatos no vienen de los árboles, que uno no se pone mostaza en los zapatos, que los cepillos de dientes no miden tres metros ni se venden en la ferretería y, que aunque las botas de piel estén hechas de piel y las botas de goma estén hechas de goma, las botas de agua no están hechas de agua. Lo único que habría que hacer para que la máquina impostora se delatara sería preguntarle sobre algunas pocas cosas de entre los tropecientos millones (esto es, de la infinidad) de asuntos humanamente obvios como éstos.

Para captar desde otro ángulo la enormidad del trabajo del programador, imagine que en el curso de la conversación con sus dos interlocutores nuestro voluntario les habla de que un hombre se toca la cabeza. ¿Cómo valorará el ordenador el posible significado de este gesto? Tocarse la cabeza con la mano puede significar que la persona tiene dolor de cabeza; que es un entrenador de béisbol que está haciendo una seña al bateador; que esa persona está tratando de ocultar su ansiedad aparentando despreocupación; que el hombre está preocupado por si se le ha deslizado el peluquín; o una infinidad de otras cosas, dependiendo de una multitud de contextos humanos en continua variación.

Existen, desde luego, lo que se conoce como sistemas expertos, programas concretos que lo hacen todo, desde analizar grandes moléculas hasta hacer cierto tipo de diagnósticos clínicos, desde escribir documentos

legales (tengo uno que redacta testamentos) hasta realizar complicados análisis estadísticos, desde recordar enormes bases de datos hasta jugar al ajedrez. Un programa clásico (si es que la frase significa algo en un campo tan reciente) llamado ELIZA incluso imita las respuestas evasivas de un psicoterapeuta no directivo, y resulta divertido lo bien que funciona durante un par de minutos.

Estos sistemas expertos están escritos por «ingenieros del conocimiento» (una expresión aterradora donde las haya). Estos programadores están adiestrados en las técnicas de la inteligencia artificial (programación designada para producir respuestas que, si procedieran de un humano, serían tenidas por inteligentes) y entrevistan a expertos en un campo determinado, pongamos geología del petróleo. Intentan captar parte de la pericia del geólogo de manera que se la pueda incorporar al ordenador: largas listas de frases acerca de rocas de la forma «Si A, B, o C, entonces D; si no E a menos que F», y complicadas redes de verdades interrelacionadas relativas a la geología. Después, si todo va bien, el sistema experto será capaz de contestar a preguntas acerca de dónde perforar para encontrar petróleo.

Es tanto más notable, pues, que con todas las cosas impresionantemente recónditas que las computadoras hacen de modo rutinario, sean precisamente la conversación, las habladurías mundanas, los chistes y las tomaduras de pelo lo que se resista más a la simulación por ordenador. (Véanse las entradas sobre *Sustituibilidad* y *La complejidad*). Simular trayectorias balísticas intercontinentales es fácil comparado con la simulación de la charla intrafamiliar. Para ésta hace falta un programa genérico multifín muchísimo más flexible. Después de escuchar a escondidas muchas conversaciones cuyos contertulios habrían pasado apuros para superar el test de Turing, quizá debería moderar un poquito mi chauvinismo humano. Sin embargo, me anima que, después de haberse enfrentado a la tarea de dotar a las máquinas de algo parecido a la inteligencia general, muchos de los que trabajan en inteligencia artificial parecen más respetuosos y están más enterados de la complejidad humana que algunos teóricos literarios. Aquéllos han tenido que tomar plena conciencia de los objetivos y fines de un programa de una manera que contrasta completamente con los esfuerzos de los desconstruccionistas, por ejemplo, por eliminarlos a ambos de sus análisis

formales y reduccionistas de los textos literarios.

Que la inteligencia artificial pueda ir más allá de los sistemas expertos para un fin concreto y cumplir su promesa (¿amenaza?) o que por el contrario acabe por ser considerada en cierto modo como un inmenso timo intelectual es algo que tardará en aclararse. Pero, si se logra la verdadera inteligencia artificial, habremos de maravillarnos de lo naturales que estas máquinas habrán llegado a ser y no de lo mecánicos que nosotros hemos sido siempre. Deberíamos pensar en nosotros mismos como pigmaliones humanos que han dado vida a sus galateas informáticas, y no como autómatas cuya base mecánica nos ha sido revelada por nuestra prole de ordenadores.

Tiempo, espacio e inmensidad

La novela del siglo XVIII de Lawrence Sterne *The Life and Opinions of Tristram Shandy, Gentleman* inspiró a Bertrand Russell su «paradoja de Tristram Shandy». La paradoja se refiere al narrador del libro, Tristram Shandy, que, según recordaba Russell, empleó dos años en escribir la historia de los dos primeros días de su vida. Shandy se lamentaba de que, a ese paso, las últimas partes de su vida nunca serían registradas. Russell señaló, no obstante, que «si hubiera vivido por siempre jamás, y no se hubiera cansado de su tarea, entonces, aun en el caso de que su vida hubiera seguido tan llena de acontecimientos como al principio, ninguna parte de su biografía hubiera quedado por escribir».

La solución de la paradoja depende de las propiedades peculiares de los números infinitos (véase la entrada sobre *Conjuntos infinitos*). Al mismo paso, Tristram Shandy habría tardado un año en escribir su tercer día, y lo mismo para los días cuarto, quinto y sexto. Cada año habría escrito el relato perfecto de otro día de su vida, y así, aunque cada año se hubiera retrasado cada vez más, no habría ningún día que quedara por registrar en el supuesto de que viviera por siempre jamás.

Quizás encontremos difícil responder racionalmente a escalas temporales muy distintas, aunque sólo nos apartemos una distancia finita de nuestros entornos contemporáneos. El intento de reconciliar los intervalos de tiempo astronómico, geológico, biológico e histórico puede producirnos un sentimiento de frustración en nanosegundos, pero a pesar de todo es una empresa que vale la pena. Si tenemos en cuenta las escalas, aparecen con frecuencia estructuras similares, y el consiguiente sentido de la perspectiva puede influir significativamente en nuestros puntos de vista, si no en nuestras

acciones y decisiones.

Con las comparaciones espaciales se tiene una perspectiva semejante. Observemos que el punto más alto de la Tierra, el Everest, solo tiene unos 9 kilómetros de altura, cifra comparable a la magnitud de la profundidad máxima de los océanos. Por tanto, las irregularidades superficiales máximas de la Tierra son menos de $1/1.000$ de sus 13.000 kilómetros de diámetro y corresponderían a abolladuras de $5/1.000$ de centímetro en la superficie de una bola de billar de 5 centímetros de diámetro (esto es, $1/1.000 \times 5$). Así pues, a pesar de las montañas, los océanos y las irregularidades del terreno, la Tierra es más lisa (aunque no necesariamente más redonda) que una bola de billar corriente.

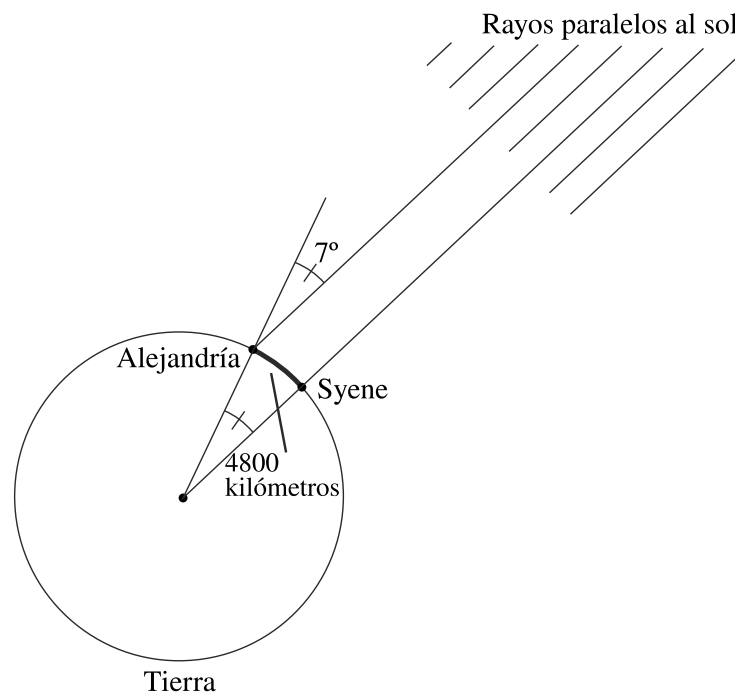
Uno de mis cuentos favoritos, «¿Cuánta tierra necesita un hombre?» de Tolstoi, no está fuera de lugar en este contexto. Trata de un hombre al que se le da la oportunidad de poseer toda la tierra que pueda rodear en un día, y la parábola muestra cómo su codicia le lleva a la muerte y así llega a la respuesta que pide el título. Un hombre necesita aproximadamente 2 metros de largo por 1 metro de ancho por 1,5 metros de profundidad: lo justo para una tumba. Pero aun siendo mucho más generosos y asignando a cada ser humano un cubículo de 7,5 metros de lado, el volumen del Gran Cañón es suficiente para alojar la totalidad de los 5.000 millones de cubos que corresponderían a los habitantes del planeta (véase la entrada sobre *Areas y volúmenes*).

Naturalmente, lo que más nos interesa son nuestros entornos espaciotemporales familiares, pero esto no habría de impedirnos ser conscientes de su necesaria limitación. Recordemos cómo la introducción en el siglo XIX de las cámaras rápidas produjo unas películas en las que la gente y los animales parecían moverse de un modo muy raro e irreal. A veces damos demasiada importancia a las divergencias relativamente poco importantes en la concepción del tiempo y los proyectos por parte de los hombres de negocios norteamericanos o japoneses, o a las pequeñas diferencias en los horizontes temporales de los adolescentes y los ciudadanos adultos. Podríamos meditar de vez en cuando sobre cómo nos relacionaríamos con un extraterrestre o con un ser artificial que, aun siendo mucho más inteligente que nosotros, tardara 100.000 veces más en responder

a los estímulos. A primera vista la comunicación con tales seres podría parecer casi imposible, sin embargo, algo vagamente parecido a la relación inversa se produce entre yo, que soy lento y listo, y mi rápido y estúpido ordenador. Y los ordenadores más rápidos son positivamente torpes comparados con los fenómenos subatómicos: un electrón da alrededor de 10^{15} vueltas al núcleo en un segundo.

A una escala mayor aún, tenemos la antigua fábula india de una piedra cúbica de una milla de lado y un millón de veces más dura que el diamante. Cada millón de años un hombre santo pasa y le hace la más suave de las caricias. Al cabo de un rato la piedra se ha gastado. La duración estimada de este rato son 10^{35} años. Para comparar, la edad del universo es de unos $1,5 \times 10^{10}$ años.

Aunque a veces me recuerde la tonta práctica masculina de pasar los primeros cinco o diez minutos de una reunión comentando el camino que se ha tomado para llegar, siempre he disfrutado estudiando las comparaciones espaciales y los símiles temporales que resumen y relacionan los distintos órdenes (astronómico, geológico, biológico, geográfico, histórico y microfísico). Aunque a veces sean simplistas, estas cartas son a pesar de todo muy útiles para orientarnos en el cosmos. El cálculo de la circunferencia de la Tierra por Eratóstenes 200 años antes de Cristo es notable en este sentido. Dedujo el valor a partir del hecho de que el sol estaba 7° al sur del cenit en Alejandría en el mismo instante en que estaba exactamente sobre la vertical de Syene,^[16] 800 kilómetros más al sur. Uno de los grandes pilares de nuestra actual concepción del mundo, la teoría de la evolución, apareció debido a la creciente insostenibilidad del marco temporal bíblico provocada por la investigación geológica. Según los estudiosos de la Biblia, que simplemente sumaban todos los «engendrados» de ésta, la edad de la Tierra era de unos 4000 años. Esta cifra tradicional se había vuelto increíble para los geólogos que estudiaban las piedras en vez de las escrituras. Después de estos descubrimientos, Darwin apareció a la vuelta de la esquina con un calendario mejor.



El sol está 7° al sur del cenit en Alejandría cuando está exactamente en el cenit en Syene, 800 kilómetros más al sur. Esto implica que el ángulo en el centro de la Tierra es también de 7° . Eratóstenes usó la proporción $C/800 \text{ kilómetros} = 360^\circ/7^\circ$ para determinar C, la circunferencia de la Tierra. C vale 40.000 kilómetros

Del conocimiento del lugar y el tiempo que uno ocupa en el mundo se deriva una cierta sensación de seguridad. Es una sensación que ha experimentado cualquier niño que escribiendo su dirección haya continuado con España, Europa, Tierra, Sistema Solar, Vía Láctea,... Despierta una sensación comparable el darse cuenta de que hemos vivido sólo 1/100.000.000 parte aproximadamente de los 4000 millones de años de la historia de la vida sobre la Tierra (suponiendo que tengamos 40 años más o menos) y que si dicha historia se comprimiera en un solo año, entonces nuestras tradiciones religiosas más «antiguas» se habrían forjado hace sólo unos 30 o 40 segundos y nosotros personalmente habríamos llegado unas 3/10 de segundo antes de la Nochevieja.

(Si podemos llegar colectivamente a las cero horas y 1 minuto del 1 de enero sin hacernos volar por los aires, me arriesgo a adelantar con audacia que estaremos tranquilos durante un buen rato).

El interés arquimedianos por el número de granos de arena que cabrían en

el universo, por mover la Tierra con una palanca muy larga, por unidades minúsculas de tiempo y de otras magnitudes cuya suma acumulada superaría necesariamente cualquier cantidad, nos hablan, todos ellos, del origen primitivo de la asociación entre la fascinación por los números y una inquietud por el tiempo y el espacio. Pascal se preguntaba por la fe, el cálculo y el lugar del hombre en la naturaleza, que está, como dijo él, a mitad de camino entre el infinito y la nada. Nietzsche especuló acerca de un universo cerrado e infinitamente recurrente. Henri Poincaré y otros, con un enfoque intuicionista y constructivista de la matemática, han comparado la sucesión de los números enteros con la concepción pre-teórica del tiempo como sucesión de instantes discretos. De Riemann y Gauss a Einstein y Gödel, los matemáticos han hecho conjeturas sobre el espacio y el tiempo. Estos temas han sido de hecho un elemento principal de la reflexión matemático-física durante milenios.

No extraigo ninguna conclusión de este discurso rudimentario, excepto quizá la de que en cierto modo esas deliberaciones nos «hacen bien»: son un tanto terapéuticas, nos hacen sentar la cabeza y tocar de pies en el suelo. Hablando de «sentar la cabeza», me molesta la gente que después de una discusión como ésta y unas cuantas copas, o bien se refugia en un dogma (no siempre religioso) o se pone sensiblero y murmura algo así como «¿Qué importancia tiene? ¿Qué importará dentro de 50.000 años?». Se podría reaccionar razonablemente con estoicismo y resignación a una pregunta fatalista como ésta. Pero piense en esto. Quizá nada de lo que hagamos ahora importe dentro de 50.000 años, pero *si* es así, entonces parecería natural que tampoco nada de lo que ocurra dentro de 50.000 años tenga importancia ahora. Y en particular no importa que dentro de 50.000 años no importe lo que hagamos ahora.

Topología

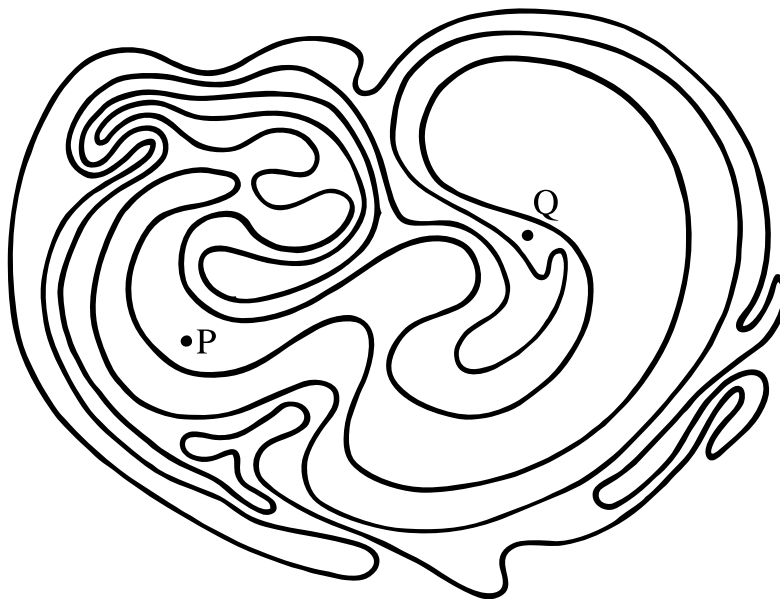
Según Woody Allen, las falsas manchas de tinta, hechas de goma, tenían originalmente un diámetro de 11 pies y no engañaban a nadie, hasta que un físico suizo «demostró que un objeto de un tamaño dado podía reducir sus dimensiones simplemente “encogiéndose”, y este descubrimiento revolucionó el negocio de las manchas de tinta de broma». Este pequeño cuento se podría interpretar como una parodia de la topología, un tema cuyas ideas pueden parecer a primera vista un tanto obvias. Se trata, con todo, de una rama de la geometría que se ocupa únicamente de aquellas propiedades básicas de las figuras geométricas que permanecen invariantes cuando las retorremos y distorsionamos, las estiramos y contraemos, o las sometemos a cualquier «deformación» siempre y cuando no las rasguemos ni las desgarramos.

En vez de dar una definición técnica de «deformación» seguiré con algunas observaciones y ejemplos más. El tamaño no es una propiedad topológica, pues como señaló el físico Woody Allen, las esferas (o las manchas de tinta de goma) se pueden hacer mayores o menores por dilatación o contracción sin necesidad de desgarrarlas, simplemente agrandándolas o encogiéndolas (pensemos en un globo que se hincha o se deshinch). Tampoco la forma es una propiedad topológica, pues un globo esférico (o una mancha de tinta con una forma rara) puede deformarse en un elipsoide, un cubo o incluso darle forma de conejo sin tener que desgarrarlo.

Como precisamente las propiedades topológicas son como las propiedades de una membrana de goma, que no cambian al estirla, comprimirla o deformarla, la topología ha sido llamada también geometría de la «membrana de goma». (Esta expresión me recuerda a mi profesor de

cálculo del instituto en Milwaukee, que fue a un cursillo de verano sobre «matemática moderna» y a partir de entonces atribuyó todas las dificultades con que se topaban sus estudiantes a la ignorancia de la geometría de la membrana de goma. Se pasaba el tiempo estirando una gran cinta de goma, como si esto ilustrara de alguna manera lo incontrovertible de su afirmación).

Si una curva cerrada en el espacio, pongamos un trozo de hilo, tiene o no un nudo es una propiedad topológica de la curva en el espacio. Una propiedad topológica de las curvas en el plano es que una curva cerrada divide el plano que la contiene en dos partes —la interior y la exterior (teorema de Jordan)—. El número de dimensiones de una figura, el hecho de tener o no borde, y en este caso de qué clase es, son también propiedades topológicas. (Véase la entrada sobre *Cintas de Möbius y orientabilidad*).

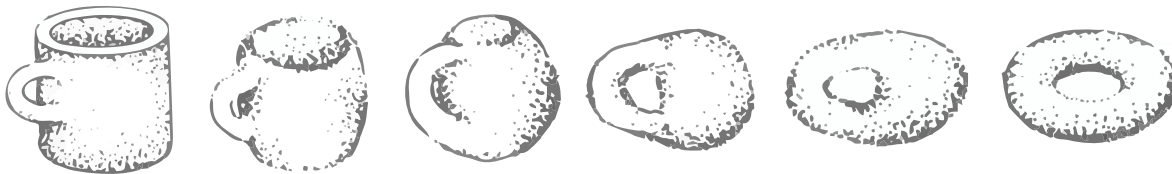


Una curva cerrada en un plano lo divide en dos partes, la interior y la exterior. ¿Está P en el interior o en el exterior?, ¿y Q?

Una cuestión que también tiene su importancia es el género de una figura: el número de agujeros que tiene o, como diría un carnicero, el máximo número de cortes que podemos hacerle sin partirla en dos trozos. Una esfera tiene género 0 pues carece de agujeros y basta un corte para romperla en dos pedazos. Un toro (una rosquilla o una figura en forma de neumático) tiene género 1, pues tiene un agujero (el agujero de la rosquilla) y se le puede hacer

un corte sin romperlo en dos pedazos. Las figuras de género 2, como unas gafas sin los cristales, tienen dos agujeros y se les pueden hacer dos cortes sin romperlas en dos partes separadas. Y así sucesivamente para las figuras con un género mayor.

Las esferas, piedras y cubos, que tienen todos ellos género 0, son topológicamente equivalentes. Para continuar con otro ejemplo de la misma índole, miremos el desayuno con los ojos de un topólogo. Henri Poincaré, uno de los fundadores de la topología (y de muchas cosas más), quizás habría observado que una rosquilla y una taza de café, ambas figuras de género 1, son topológicamente equivalentes. Para verlo, imaginemos una taza de café hecha con arcilla. Aplanemos el cuerpo de la taza y agrandemos el tamaño del asa, estrujando la arcilla para que pase del uno a la otra. El agujero del asa se convierte así en el agujero de la rosquilla y podemos percibir fácilmente la equivalencia topológica. Los seres humanos tienen también, al menos *grosso modo*, género 1. Somos topológicamente equivalentes a las rosquillas, y nuestro canal digestivo/excretor correspondería al agujero de la rosquilla. (Pero esto último tiene menos encanto que la inútil especulación del desayuno).



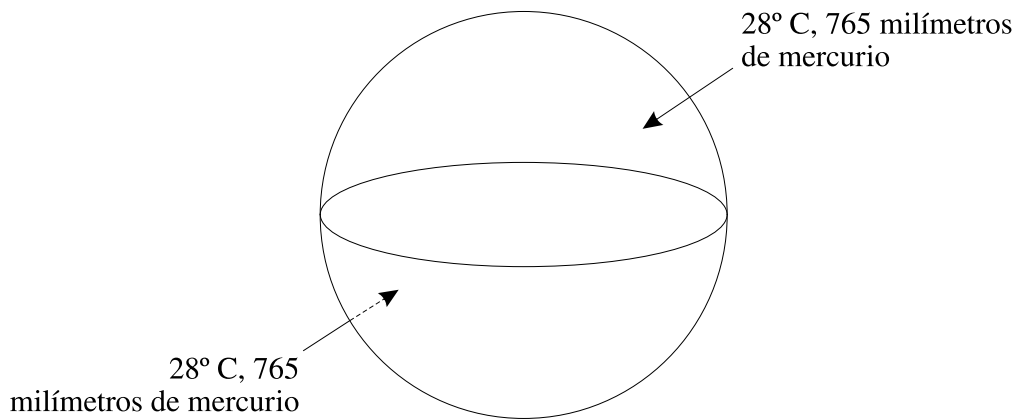
Una taza de café de arcilla con una asa se puede deformar continuamente hasta convertirla en una rosquilla. Son topológicamente equivalentes.

Estas ideas tienen varias aplicaciones, pero en su mayoría son internas a la propia matemática. Frecuentemente en el trabajo teórico, por ejemplo, es importante saber que existe una solución pero no hace falta tener un método para encontrarla. Para hacerse una idea de esos llamados teoremas de existencia, pensemos en un escalador que empieza su ascenso a las seis de la mañana de un lunes y llega a la cumbre a mediodía. Y empieza el descenso el martes por la mañana a las seis y llega al pie a mediodía. No suponemos nada más acerca de si va rápido o lento en su ascenso y descenso de estos dos días.

Podría, por ejemplo, haber subido a un ritmo lento, descansando a menudo, en su ascenso del lunes, y, después de pasear despreocupadamente cerca de la cumbre el martes por la mañana, haber bajado literalmente a mata caballo, incluso cayendo los últimos 300 metros. La pregunta es: ¿podemos estar seguros de que, independientemente de cómo suba o baje, habrá necesariamente un instante entre las seis de la mañana y el mediodía en el cual el escalador estará exactamente a la misma altura, tanto a la subida como a la bajada?

La respuesta es sí, y la demostración es clara y convincente. Imaginemos el ascenso y el descenso, en todos sus detalles, realizados simultáneamente por dos escaladores. Uno empieza en el pie y el otro en la cumbre y ambos parten a las seis de la mañana, imitando lo que el escalador original hizo el lunes y el martes, respectivamente. Está claro que estos escaladores se cruzarán en algún punto del camino y que en este instante los dos estarán a la misma altura. Como no hacen sino reconstruir los pasos del escalador original, podemos concluir con toda seguridad que éste estaba a la misma altura a la misma hora de los dos días sucesivos.

Un ejemplo menos intuitivo de teorema de existencia es el resultado de que siempre hay un par de puntos diametralmente opuestos (antipodales) sobre la superficie de la Tierra que tienen la misma temperatura y la misma presión barométrica. Estos puntos van variando y no tenemos manera de encontrarlos, pero podemos demostrar que existen siempre. No se trata de un fenómeno meteorológico, sino matemático. Otro ejemplo: tomemos un pedazo de papel rectangular y coloquémoslo plano en el fondo de una caja, poniendo cuidado en que todo el fondo quede perfectamente cubierto. Si ahora arrugamos el papel y hacemos una pelota con él y lo dejamos en la caja, podemos estar topológicamente seguros de que al menos un punto del papel está precisamente en la vertical del mismo punto del fondo de la caja que cubría antes de arrugar el papel. La existencia de dicho punto fijo es segura.



Siempre hay un par de puntos antipodales que tienen exactamente la misma temperatura y presión

Teoremas como éstos a veces nos llevan a resultados concretos y prácticos en campos como las teorías de grafos y redes, que se ocupan, entre otras cosas, de idealizaciones matemáticas de las redes de calles y autopistas. Pero, como ya dije, contribuyen más frecuentemente a avances teóricos en otras ramas de la matemática. La topología algebraica, por ejemplo, usa ideas topológicas y algebraicas para caracterizar diversas estructuras geométricas, mientras que la topología diferencial emplea técnicas de las ecuaciones diferenciales y de la topología para estudiar tipos muy generales de variedades (superficies) de muchas dimensiones. Y la teoría de catástrofes, una subdisciplina de la topología diferencial, se ocupa de la descripción y clasificación de discontinuidades —pliegue, cúspide, mariposa, ...—. La topología es mucho más que las manchas de tinta de broma.

El triángulo de Pascal

El triángulo de Pascal es una disposición triangular de números que ya conocían los chinos en 1303, más de 300 años antes de que el matemático y escritor francés Blaise Pascal descubriera muchas de sus propiedades más interesantes. Exceptuando los unos de los lados, cada número del triángulo se obtiene sumando los dos números que están encima suyo. Así, si hubiéramos de añadir una nueva fila a la figura de más abajo, tendríamos 1, 7, 21, 35, 35, 21, 7, 1.

				1		1			
			1		2		1		
		1		3		3		1	
	1		4		6		4		1
1		5		10		10		5	1
1	6		15		20		15	6	1

Triángulo de Pascal

Aunque la regla anterior es bastante simple, es sorprendente la diversidad de pautas contenidas en el triángulo. Probablemente la más importante de ellas tenga que ver con los números combinatorios (o coeficientes binomiales). Estos números nos dicen, entre otras cosas, cuantas manos de póquer, billetes de lotería y cuantos helados de tres gustos de Baskin-Robbins posibles hay. Y, en general, nos dan el número de modos distintos en que podemos elegir R elementos de un conjunto de N .

Para ilustrarlo con un número pequeño, consideremos la cuarta fila del triángulo: 1, 4, 6, 4, 1. Estos cinco números nos dan el número de modos en que podemos escoger 0, 1, 2, 3 y 4 elementos, respectivamente, de un conjunto de cuatro elementos. Así, el primer número de la fila, 1, indica el número de maneras de elegir 0 elementos de un conjunto de 4. Sólo hay un modo de hacerlo: no tomar ninguno. El siguiente número, 4, indica el número de maneras de elegir 1 de los 4 elementos. Astutamente, observamos que tenemos cuatro posibilidades ahora: elegir el primer elemento, el segundo, el tercero o el cuarto. El tercer número de la fila, 6, nos indica los modos en que podemos elegir 2 elementos de los cuatro. Para ilustrarlo mejor, supongamos que los elementos son las letras A, B, C, D. Las seis maneras de tomar exactamente 2 letras son: AB, AC, AD, BC, BD y CD. El siguiente número de la fila, 4, es el número de modos en los que podemos escoger 3 de los cuatro elementos. Hay cuatro maneras de hacerlo: basta con decir cuál de los cuatro no tomamos. Y el número siguiente es 1, pues sólo hay un modo de elegir 4 elementos de 4: tomarlos todos.

Todas las filas funcionan igual. Los números de la sexta fila —1, 6, 15, 20, 15, 6, 1— representan, respectivamente, el número de modos en que podemos tomar 0, 1, 2, 3, 4, 5 y 6 elementos de un conjunto de 6. Vemos pues que hay 15 maneras de escoger 2 elementos de entre 6 y 20 maneras de tomar 3 elementos. Si miramos la fila 31.^a del triángulo de Pascal, en el cuarto lugar encontramos el número de maneras posibles de escoger 3 sabores de entre los 31 que ofrecen los helados Baskin-Robbins: 4.495. Si buscamos en la fila 49 y contamos 7 lugares, encontramos el número de modos de elegir 6 elementos de entre 49: esto es, el número de combinaciones posibles en la Lotería Primitiva: 13.983.816. Si vamos a la fila 52.^a y contamos 6 lugares, encontramos el número de maneras de escoger 5 elementos de entre 52, esto es, el número de manos de póquer posibles: 2.598.960.

(Me siento tentado a dar más ejemplos, pero siempre que llego a una situación así, en que cito un ejemplo tras otro, todos ellos muy parecidos, el recuerdo de mi primer profesor de cálculo en la universidad hace que controle mi inclinación a la pedantería y me detenga. Este profesor tenía 200 alumnos en clase, en un aula muy grande. Un buen día empezó a comportarse de

manera un tanto extraña, pero, como eso no era poco habitual en él, no le hicimos mucho caso. Estaba hablando de series condicionalmente convergentes y se puso a escribir algo así como $1 - 1/2 + 1/3 - 1/4 + 1/5 - \dots$ en el lado izquierdo de la pizarra. Siguió así y cuando llegó a la mitad del encerado, que era muy ancho, andaba por el $1/57 - 1/58 + \dots$. Tomándolo por otra muestra de sus excentricidades, la clase se divertía con ello, pero cuando se acercaba a la derecha y seguía $- 1/124 + 1/125 - \dots$, nos fuimos quedando callados. Por fin llegó al final de la pizarra, se volvió y nos miró. La mano con la que sostenía la tiza tembló por un instante, la dejó caer y a continuación abandonó el aula. No volvió más a clase y más tarde nos contaron que había sufrido una crisis nerviosa).

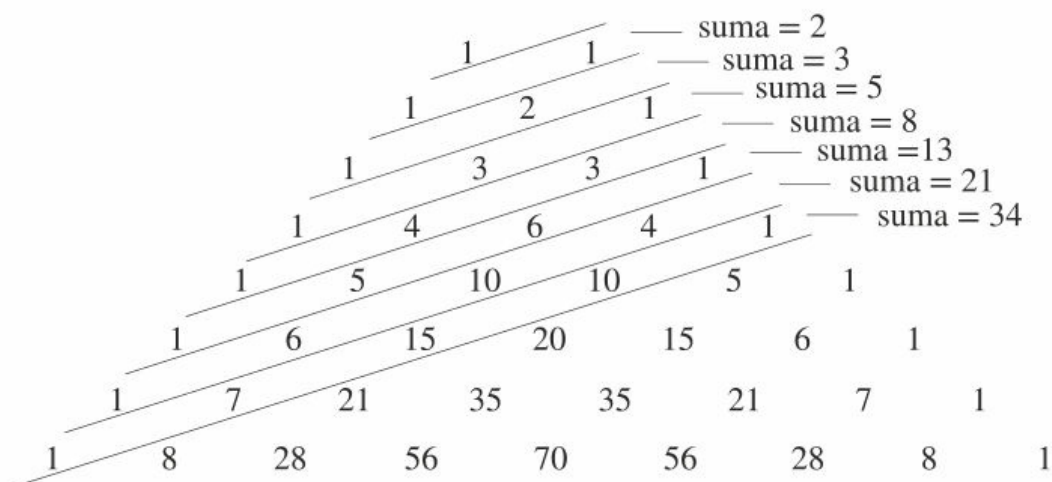
Para números grandes esta manera de obtener los llamados números combinatorios no es muy práctica, y por ello se usa una fórmula obtenida a partir de la regla del producto (véase la entrada sobre *La regla del producto*). Dice que el número de maneras de escoger R elementos de un conjunto de N , que se representa normalmente por $C(N, R)$, es $N!/R!(N - R)!$, donde, para cualquier número X , $X!$ indica el producto de X por $X - 1$, por $X - 2$, y así hasta 1. Por ejemplo, $6! = 6 \times 5 \times 4 \times 3 \times 2 \times 1 = 720$. Si comprobamos que la fórmula da el mismo resultado que el triángulo de Pascal para $N = 6$ y $R = 2$, observamos que $6!/2!4! = 15$, el número de maneras de tomar 2 elementos de un conjunto de 6; $C(6,2) = 15$.

Esta fórmula es particularmente valiosa en la teoría de la probabilidad y en combinatoria, donde se suele recurrir al recuento de posibles situaciones. Por ejemplo, si uno ignora una materia por completo pero ha de realizar un test sobre la misma con 12 preguntas de 5 respuestas posibles cada una, la probabilidad de que conteste bien a las 4 primeras y mal a las 8 restantes es $(1/5)^4 \times (4/5)^8$. Ésta es también la probabilidad de que conteste bien a las preguntas 3, 4, 7 y 11 y mal al resto. Para determinar la probabilidad de contestar bien a cualquier conjunto de exactamente 4 preguntas, encontramos primero el número de maneras distintas de escoger 4 de entre 12, esto es $C(12,4)$ o 495, y multiplicamos este número por $(1/5)^4 \times (4/5)^8$, la probabilidad de contestar bien a un conjunto particular de 4 preguntas. El resultado, $C(12,4) \times (1/5)^4 \times (4/5)^8$, es la respuesta (aproximadamente el

13%) y también un caso especial de la importante distribución binomial en teoría de la probabilidad.

Hay que decir también que los números de la N-ésima fila del triángulo de Pascal son los coeficientes del desarrollo de $(X + Y)^N$. Para ilustrar esto, recuerden (o creen en mi palabra) que $(X + Y)^2$ es igual a $1X^2 + 2XY + 1Y^2$ y que $(X + Y)^3$ es igual a $1X^3 + 3X^2Y + 3XY^2 + 1Y^3$ y fíjense en que los números subrayados coinciden con las filas 2.^a y 3.^a del triángulo de Pascal. Y también $(X + Y)^4 = 1X^4 + 4X^3Y + 6X^2Y^2 + 4XY^3 + 1Y^4$.

Algunas otras pautas agazapadas en el triángulo de Pascal son: los números enteros a lo largo de las penúltimas diagonales, los números triangulares (expresables por disposiciones triangulares de puntos –3, 6, 10, etc.) a lo largo de las diagonales siguientes según se baja hacia el centro, los números tetraédricos (expresables por disposiciones tetraédricas de puntos –4, 10, 20, etc.) sobre las diagonales siguientes, los análogos a ellos para dimensiones superiores conforme se toman diagonales más interiores, y los números de Fibonacci como sumas de los elementos de las diagonales que suben (bajan) de una fila a la anterior (siguiente). El hallazgo de estas configuraciones, y de otras semejantes, nos hace apreciar la belleza y la complejidad que a veces es inherente a la más simple de las reglas.



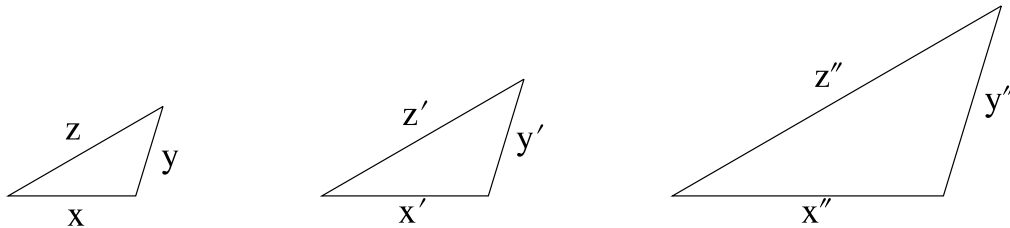
Los números de Fibonacci se obtienen del triángulo de Pascal

Trigonometría

La trigonometría data de los antiguos griegos y Eratóstenes ya la empleó para calcular la circunferencia de la Tierra, Aristarco para estimar las distancias al Sol y la Luna, y muchos otros para perforar túneles en línea recta y calcular distancias a puntos inaccesibles. Durante siglos los astrólogos la usaron para construir sus cartas y para sus cálculos, y aún hoy muchos explotan la imagen esotérica de la trigonometría celeste para disfrazar la falta de fundamento de sus creencias astrológicas. La parte elemental del tema, en la que nos centraremos principalmente aquí, trata de los triángulos rectángulos y de los cocientes entre sus distintos lados, temas que quizás a primera vista no parezcan nada atractivos, y tampoco a segunda. Sin embargo, está en los programas de estudios (normalmente en un lugar demasiado importante, en mi opinión), así que perseveremos y consideremos un ejemplo canónico. Supongamos que nos encontramos a 30 metros del pie de una torre repetidora de televisión cuya altura queremos conocer. Supongamos también que el ángulo de elevación de la torre, el ángulo que forman nuestros ojos con el suelo cuando miramos a la cima, es de 45° . ¿Cuál es la altura de la torre?

La idea básica de la trigonometría elemental es que para triángulos semejantes, aquellos que tienen los lados proporcionales, si determinamos la razón entre un par de lados de uno de ellos, encontraremos que el par de lados correspondientes del otro triángulo guardan la misma proporción. Así, en triángulos rectángulos cuyos ángulos no rectos midan 45° cada uno (recordemos que la suma de los tres ángulos de un triángulo es siempre 180° y que los triángulos semejantes tienen los ángulos iguales), la razón entre los dos lados más cortos es siempre de 1 a 1, independientemente de lo grandes o

pequeños que sean los triángulos. Como estamos a 30 metros del pie de la torre, deducimos que la altura de esa torre de televisión es de 30 metros.



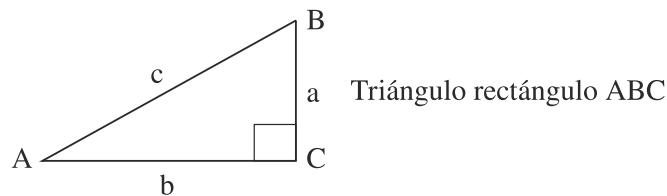
En los triángulos semejantes los lados correspondientes son proporcionales

$$\frac{x}{y} = \frac{x'}{y'} = \frac{x''}{y''}; \frac{y}{z} = \frac{y'}{z'} = \frac{y''}{z''}; \frac{z}{x} = \frac{z'}{x'} = \frac{z''}{x''}$$

En general, la tangente de un ángulo agudo (menos de 90°) de un triángulo rectángulo se define como el cociente del lado opuesto al ángulo entre el lado adyacente a él. En el ejemplo de la torre de televisión la tangente del ángulo de elevación es la razón de la altura de la torre a nuestra distancia a ella, y cuando este ángulo es de 45° , la tangente es 1 entre 1, que puesto en forma de cociente es 1/1 o simplemente 1. Si estuviéramos a 30 metros de la torre, pero el ángulo de elevación fuera sólo de 20° , entonces podríamos consultar una tabla trigonométrica para determinar que la tangente de este ángulo es aproximadamente 0,36 entre 1, o simplemente 0,36. En este caso, la altura de la torre de televisión sería sólo de $0,36 \times 30$, es decir, unos 10,8 metros.

La tangente de un ángulo es sólo una de entre las varias funciones trigonométricas, de nombres un tanto soporíferos (seno, coseno, etc.), que asocian cocientes y números con ángulos. (Conocí a alguien que para referirse al seno de X decía Sominex).^[17] Con frecuencia conocemos algunos ángulos y algunas longitudes, y por medio de estas funciones podemos determinar los ángulos y longitudes desconocidos que queremos encontrar. El seno de un ángulo agudo en un triángulo rectángulo se define como la razón entre el lado opuesto al ángulo y la hipotenusa (el lado opuesto al ángulo recto), mientras que el coseno de un ángulo es la razón entre el lado adyacente al ángulo y la hipotenusa. En el caso del ángulo de 45° , el seno y el

coseno son iguales y valen aproximadamente 0,71, mientras que para un ángulo de 20° , el seno es 0,34 y el coseno 0,94. Simbólicamente escribimos: $\text{sen}(20^\circ) = 0,34$ y $\text{cos}(20^\circ) = 0,94$, mientras que $\text{tg}(20^\circ) = 0,36$ y $\text{tg}(45^\circ) = 1$.



La tangente del ángulo A es $\frac{a}{b}$; $\text{tg}(A) = \frac{a}{b}$

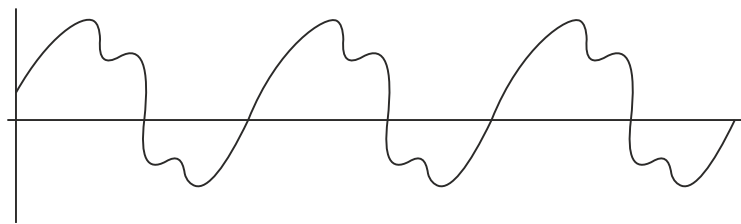
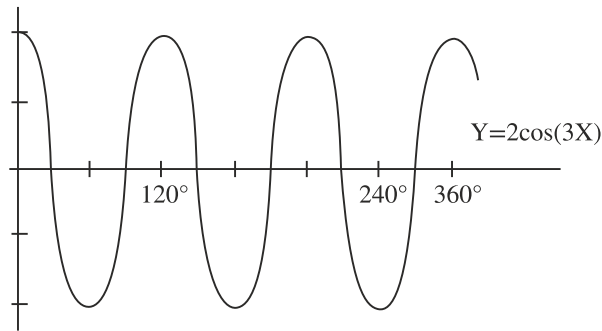
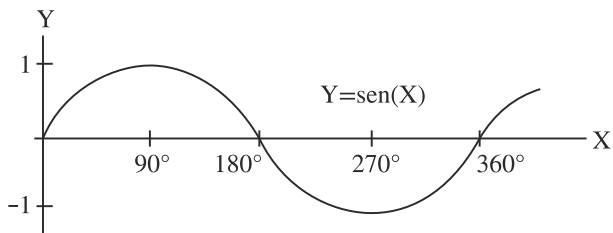
El seno del ángulo A es $\frac{a}{c}$; $\text{sen}(A) = \frac{a}{c}$

El coseno del ángulo A es $\frac{b}{c}$; $\text{cos}(A) = \frac{b}{c}$

Funciones trigonométricas más comunes.

Los valores de las funciones trigonométricas de cualquier ángulo se pueden encontrar consultando las tablas compiladas a tal efecto, usando las fórmulas de series infinitas (véase la entrada sobre *Series*) o, más fácilmente, presionando las teclas apropiadas de una calculadora. En la Antigüedad y en la Edad Media no había ni tablas, ni fórmulas, ni calculadoras y tenían que recurrir a métodos geométricos complicados para encontrar el seno, el coseno o la tangente de los ángulos que les interesaban, pero su modo de razonar no era muy distinto del que hemos usado en nuestro ejemplo de la torre de televisión o en los actuales problemas de agrimensura, navegación o astronomía.

Habría de intercalar aquí que, a pesar de lo que crean muchos estudiantes de primer curso, el seno de un ángulo de 60° no es el doble del seno de un ángulo de 30° , ni el triple del seno de un ángulo de 20° . Idem para los cosenos y las tangentes. Las funciones trigonométricas no crecen de este modo proporcional. Además, como son funciones y no productos, simplificar el 20° en $\text{sen}(20^\circ)/20^\circ = \text{sen}$ es efectivamente un pecado mortal matemático, pero además, matemáticamente hablando, carece de sentido. (Véase la entrada sobre *Funciones*).



Suma de varias ondas sinusoidales distintas.

Como suele ocurrir, el centro de interés de una disciplina va cambiando y las definiciones se generalizan. En el caso de la trigonometría se hizo necesario ampliar las definiciones de las funciones trigonométricas para tratar con triángulos no rectángulos y extenderlas al caso de los ángulos obtusos, de más de 90° . Estas definiciones más generales se formulan en términos de círculos y rotaciones, y relacionan la trigonometría elemental con sus materializaciones más modernas.

El seno de un ángulo varía. Vale 0 para un ángulo de 0° , crece continuamente, pero no linealmente, hasta alcanzar un valor máximo de 1 para un ángulo de 90° , vuelve a disminuir hasta 0 para 180° , se hace más y más negativo, bajando hasta -1 para 270° , aumenta gradualmente hasta volver a 0 para 360° , y vuelve a empezar este mismo ciclo para ángulos mayores de 360° . (Un ángulo de 370° es indistinguible de uno de 10°). Al representarla gráficamente, esta variación periódica del seno de un ángulo

produce la típica forma u onda sinusoidal, que describe muchos fenómenos físicos y, en particular, la corriente eléctrica alterna.

La trigonometría moderna se interesa más por la periodicidad y otras propiedades de las funciones trigonométricas que por su interpretación en términos de cocientes. En su estudio inicial del calor, el matemático francés Joseph Fourier (chiste matemático estúpido: el nombre se pronuncia *ye ye ye* exactamente igual que *tran tran tran tran* es el nombre de un conocido lenguaje informático)^[18] sumó series de senos y cosenos de distintas frecuencias (grados de movimiento) para imitar el comportamiento de las funciones periódicas no trigonométricas. A modo de analogía, imaginemos que combinamos los sonidos procedentes de dos diapasones que vibran a frecuencias distintas para obtener un sonido que es la «suma» de los dos. Estas técnicas de suma nos permiten aproximar una gran familia de funciones importantes, incluso las no periódicas, con series de Fourier infinitas de funciones trigonométricas, y casi dos siglos de trabajo en esas series de Fourier las han convertido en un instrumento de cálculo indispensable para la ciencia y la ingeniería, y han producido al mismo tiempo teoremas sutiles y profundos de matemática pura.

Sin embargo, mucha gente recuerda las identidades trigonométricas, suponiendo que recuerden algo de ellas, como fórmulas que indican la igualdad incondicional entre expresiones trigonométricas complicadísimas. Frecuentemente se emplea una barbaridad de tiempo de la clase de trigonometría en realizar las manipulaciones formales necesarias para demostrar dichas identidades. De hecho, sólo unas pocas de ellas son cruciales y, probablemente, la más importante de ellas sea la que dice que el cuadrado del seno de un ángulo más el cuadrado del coseno del mismo ángulo siempre da igual a 1. Si lo comprobamos con los valores que hemos dado más arriba, vemos que para el ángulo de 45° , $0,71^2 + 0,71^2 = 1$, y para el de 20° , $0,34^2 + 0,94^2 = 1$.

Acabaré con un pequeño rompecabezas y una historieta. La urgencia del primero ha ido disminuyendo con la generalización del uso de los relojes digitales. Usted tiene un reloj de pulsera cuyas manecillas se mueven de manera continua y observa que son las tres. ¿Cuánto tardará la minutería en

alcanzar a la horaria? La respuesta aproximada es obvia, entre 15 y 20 minutos, pero ¿cuál es la respuesta exacta? Y ahora la historietta. Recientemente estaba escuchando en la radio el programa de un psicólogo. Su invitado, un famoso de segundo orden, se puso a hablar soporíferamente sobre cómo había llegado a tocar fondo por culpa de las drogas y el alcohol, pero que, gracias al desarrollo espiritual y a sus esfuerzos en una clínica de rehabilitación, recientemente había «conseguido dar a su vida un giro de 360° ».

[La solución al rompecabezas: Consideremos el ángulo que se ha movido la horaria cuando la alcanza la minuteria, llamémosle X . Como en este tiempo la minuteria se mueve 12 veces más rápido (da una vuelta completa mientras la horaria va de las 3 a las 4, $1/12$ de vuelta), el ángulo recorrido por la minuteria es $12X$. Pero otra manera de caracterizar este ángulo es que la minuteria se ha movido $(90 + X)$, el ángulo que se ha movido la horaria más los 90° de desventaja que llevaba la primera respecto a la segunda, pues aquélla empezó apuntando a las 12 y ésta a las 3. De la igualdad de ambas expresiones tenemos: $12X = (90 + X)$. Despejando obtenemos: $X = 8,1818^\circ$. Y traduciéndolo a minutos y segundos ($90^\circ = 15$ minutos), encontramos que la minuteria alcanza a la horaria a las 3:16:22].

Variables y pronombres

Una variable es una cantidad que puede tomar distintos valores, pero cuyo valor en una situación dada es a menudo desconocido. Es lo contrario de una cantidad constante. El número de padres biológicos de una persona es una cantidad constante. El número de sus retoños es una variable.

Sorprendentemente, no fue hasta finales del siglo XVI cuando al matemático francés François Viète se le ocurrió la idea, que retrospectivamente parece obvia, de usar letras para representar las variables (normalmente X , Y y Z para los números reales y N para los enteros). A pesar de las protestas de generaciones de estudiantes principiantes en álgebra por la introducción de las variables, su uso no es más abstracto que el de los pronombres, con los que guardan un fuerte parecido conceptual. (Los nombres, por contra, son los análogos de las constantes). Y al igual que los pronombres hacen la comunicación más fácil y más flexible, las variables nos permiten trabajar con una mayor generalidad que si limitamos nuestro discurso matemático a las constantes.

Consideremos la frase siguiente. «En cierta ocasión alguien dio a su mujer algo que ella encontró tan desagradable que lo tiró al cubo de basura más próximo y nunca volvió a mencionarlo de buena gana, a pesar de que de vez en cuando él le preguntaba por su paradero». Sin pronombres la misma frase sería muy farragosa: «En cierta ocasión esta persona dio a la mujer de esta misma persona una cosa, y la mujer de esta persona encontró esta cosa tan desagradable que la mujer de esta persona tiró esta cosa al cubo de basura más próximo y nunca volvió a mencionar de buena gana esta cosa a esta persona, a pesar de que esta persona preguntaba de vez en cuando por el paradero de esta cosa a la mujer de esta persona». Si introducimos variables

la frase recupera un poco su manejabilidad: «X dio a Y, mujer de X, un Z e Y encontró Z tan desagradable que tiró Z al cubo de basura más próximo y nunca volvió a mencionar de buena gana Z a X, aunque X de vez en cuando preguntaba a Y por el paradero de Z».

Tenemos un ejemplo breve en el mandato a Oscar: «Ayuda a quien te ayude». Sin pronombres debería decir: «Ayuda a Jorge si Jorge ayuda a Oscar, ayuda a Pedro si Pedro ayuda a Oscar, ayuda a Marta si Marta ayuda a Oscar, ayuda a Juana si Juana ayuda a Oscar, etc.».

Dado que el uso de los pronombres y todo lo relacionado con ellos no representa un problema para casi nadie, parece pues que poca gente habría de tener dificultades con las variables. Sin embargo, en matemáticas se imponen condiciones a las variables que frecuentemente nos permiten determinar su valor. Si $X - Y + 2(1 + 3X) = 31$ e $Y = 3$, podemos encontrar X. Son las técnicas que se emplean para resolver estas ecuaciones y otras más complicadas lo que a menudo resulta un enigma. En nuestro discurso cotidiano con los pronombres no hay ninguna situación cuya analogía con el caso matemático sea obvia, pero las novelas de misterio no son tan distintas del mismo como pudiera pensarse. Consideremos el siguiente ejemplo: quienquiera (el señor X o la señora X) que anulara las reservas de hotel de los invitados sabía que venían a la celebración, que llegarían tarde y que si no tenían una reserva a su nombre les causaría molestias, a ellos y a sus huéspedes. Si conocemos los principales personajes implicados en ello, ¿podemos descubrir quien anuló las reservas (esto es, quién es igual a X)? Estoy convencido de que las técnicas y aproximaciones que se emplean para aclarar y resolver pequeños dramas humanos como éstos son por lo menos tan complejos como las que se usan en matemáticas. (Véase también la entrada sobre *Sustituibilidad*).

Un último comentario editorial: algunos han argumentado que la naturaleza teórica de las matemáticas nos aleja de nuestra humanidad y es, en cierto modo, incompatible con el espíritu de compasión. Sin embargo, como ya he sugerido en esta entrada y en otros lugares, el lenguaje que empleamos corrientemente contiene toda la abstracción de la matemática. El «problema» de ésta no es que sea abstracta, sino que demasiado a menudo su abstracción está poco fundamentada, sin una base lógica humana. En cuestiones de

política social o toma de decisiones personales las matemáticas pueden servir para determinar las consecuencias de nuestras hipótesis y valores, pero el origen de éstos y aquéllas está en nosotros (nosotros X), y no en unas divinidades matemáticas.

Zenón y el movimiento

Hacia el año 460 a. C., unos ochenta años después de Pitágoras, pensaba y escribía el filósofo griego Zenón de Elea. Aunque sus obras no nos han llegado, las referencias de Aristóteles nos sugieren un agudo escéptico cuyas diversas paradojas no pudieron ser resueltas por los matemáticos de la época. La más famosa de estas paradojas es la de la carrera entre Aquiles y la tortuga y parece demostrar que, después de haber dado a la tortuga una ventaja de salida, Aquiles nunca puede alcanzarla, independientemente de lo deprisa que vaya. Para alcanzar la tortuga, Aquiles debe llegar primero al punto de partida de ésta, T_1 . Pero en este tiempo la tortuga habrá avanzado hasta el punto T_2 , y Aquiles habrá de apresurarse a cubrir la distancia entre T_1 y T_2 . Pero mientras, la tortuga habrá avanzado hasta T_3 , y mientras Aquiles va de T_2 a T_3 , la tortuga se habrá movido hasta T_4 .

Así pues, razonaba Zenón, Aquiles nunca adelantará a la tortuga, porque para hacerlo tendría que realizar un número infinito de actos en un período de tiempo finito. Esto es, Aquiles ha de recorrer la distancia entre T_0 y T_1 , entre T_1 y T_2 , entre T_2 y T_3 , entre T_3 y T_4 , entre $T_{17.385}$ y $T_{17.386}$, etc. Y como para recorrer cada una de estas distancias se tarda un cierto tiempo, para recorrer una infinidad de ellas se necesitará un tiempo infinito. Por tanto, Aquiles nunca alcanzará a la tortuga, aunque cada vez se le acercará más. La conclusión es obviamente falsa, pero ¿dónde está el fallo del argumento?



Para alcanzar a la tortuga Aquiles ha de atravesar la distancia entre T_0 y T_1 entre T_1 y T_2 , entre T_2 y T_3 , etc. Zenón razonaba que, como para atravesar cada una de estas distancias hace falta algún tiempo, para atravesar una infinidad de ellas haría falta un tiempo infinito

Antes de describir la solución al dilema, presentaré un par más de ellos. En la paradoja de la flecha, Zenón sostenía que una flecha está quieta incluso cuando está en mitad de su vuelo. En un instante particular cualquiera, la flecha simplemente está donde está y ocupa un volumen de espacio exactamente igual al suyo propio. En este instante, durante este momento preciso, la flecha no puede moverse pues, de hacerlo, se seguiría una de las dos consecuencias absurdas siguientes. El movimiento implicaría que el instante tiene una parte inicial y una parte final, pero por definición los instantes no tienen partes. O si no, implicaría que la flecha ocupa un volumen de espacio mayor que el suyo, para tener espacio para poder cambiar de posición. Ninguna de las dos posibilidades tiene sentido y, por tanto, deducimos la imposibilidad del movimiento durante este instante. Concluimos también que el movimiento es imposible porque, de producirse, habría de tener lugar durante un instante u otro.

Como antes, la conclusión es claramente falsa, pero ¿qué falla? El genio de Zenón consistió precisamente en su disposición a seguir los argumentos hasta sus últimas consecuencias, aunque le llevaran a posiciones contradictorias, en este caso relativas al espacio y al tiempo. ¿Son infinitamente divisibles el espacio y el tiempo? ¿Cómo se explica entonces la paradoja de Aquiles y la tortuga? ¿Son discretos y espasmódicamente cinematográficos? ¿Cómo se explica entonces la paradoja de la flecha?

Y un último enigma que, aunque no es comparable a los anteriores por su importancia histórica sino que más bien es un truco, también va de movimiento y deporte griego, y quizás éste sea su lugar apropiado. Consideremos dos corredores de la categoría sénior que corren una maratón —la distancia aproximada es de unos 42,2 kilómetros—. Uno de los corredores, Jorge, marcha a un ritmo constante de 6 minutos por kilómetro. El otro, Juan, va a un ritmo bastante desigual, pero se sabe que si le cronometramos cualquier tramo de un kilómetro de longitud de la carrera tarda siempre 6 minutos y 1 segundo. La pregunta es: ¿puede Juan ganar a

Jorge en la maratón a pesar de su ritmo persistentemente más lento?

La respuesta es que sí, naturalmente ¿Por qué, si no, habría planteado el problema? Podría suceder que Juan «esprintara» los primeros 200 metros en 25 segundos y luego recorriera los siguientes 800 metros en 5 minutos y 36 segundos. Y fuera repitiendo esta pauta alternante de *sprints* de 25 segundos seguidos de recuperaciones de 5 minutos y 36 segundos. Como la maratón son 42,2 kilómetros, el tiempo total de Jorge será 253,2 minutos (42,2 kilómetros $\times$ 6 minutos/kilómetro) o 253 minutos y 12 segundos; el tiempo total de Juan será 253 minutos y 7 segundos (42 kilómetros $\times$ 6 minutos 1 segundo/kilómetro más los 25 segundos de los últimos 200 metros). Así pues, gana Juan por 5 segundos.

¿Cuáles son las soluciones modernas a las paradojas de Zenón? En la de Aquiles, Zenón supone incorrectamente que la suma total de una infinidad de intervalos (en los cuales Aquiles va de T_0 a T_1 , de T_1 a T_2 , etc.) es infinita, de lo que deduce que Aquiles nunca alcanza a la tortuga. Que no tiene por qué ser necesariamente así se ve considerando la serie $1 + 1/2 + 1/4 + \dots + 1/8 + 1/16 + 1/32 + \dots$, que, a pesar de tener una infinidad de términos, sólo suma 2. Estos temas no se aclararon completamente hasta el siglo XIX, con la rigorización del cálculo y las series infinitas (véanse también las entradas sobre *Series* y *Límites*).

En cuanto a la paradoja de la flecha, Zenón tenía razón al pensar que en un instante particular la flecha está en una posición particular. También estaba en lo cierto al creer que no hay ninguna diferencia intrínseca entre que una flecha esté en reposo en un instante dado y que esté en movimiento en este mismo instante; el movimiento y el reposo son instantáneamente indistinguibles. Su error fue deducir de ello que el movimiento era imposible. La diferencia entre el reposo y el movimiento se manifiesta sólo cuando consideramos las posiciones de la flecha en distintos instantes de tiempo. El movimiento consiste precisamente en estar en lugares distintos en instantes distintos, y el reposo en estar en el mismo lugar en instantes distintos.

Un aspecto notable de estas paradojas es lo que se tardó en resolverlas. Se necesitan algunas de las técnicas más sutiles del cálculo y las series para aclarar los enigmas planteados hace 2.500 años por Zenón de Elea. Estos experimentos teóricos nos proporcionan pues un magnífico ejemplo de ideas

y narrativa matemáticas anteriores a las ecuaciones y los cálculos.

Éste es el orden natural de desarrollo, pero se invierte con demasiada frecuencia, especialmente en la pedagogía matemática. Cuando esto pasa, las matemáticas se convierten en una colección de técnicas, y se pierde su relación íntima con la filosofía, la literatura, la historia, la ciencia y la vida cotidiana. Sin el apoyo de su contenido humano, la matemática deja de ser una de las artes liberales para convertirse en un simple instrumento técnico. Creo que hay que oponerse a ello y he aquí mi contribución a la resistencia activa.

Lista cronológica de «Los cuarenta principales»

A mucha gente le gusta ver una lista de los «cuarenta principales» a pesar de las distorsiones, errores y tonterías que inevitablemente produce algo tan simplista. Cuando se trata de discos hay al menos un patrón más o menos claro en base al cual elaborar la lista: el número de discos vendidos en un cierto intervalo de tiempo. Para matemáticos distribuidos a lo largo de varios milenios no tenemos un patrón semejante, que tenga en cuenta conjuntamente el número de teoremas demostrados, el número de citas, la importancia y profundidad de sus descubrimientos, más una buena cantidad de *je ne sais quoi*. No obstante, la que sigue es una lista de matemáticos (la mayoría de ellos ya han sido citados en el libro) que generalmente han sido considerados entre los «mejores» en un período que abarca desde la antigüedad hasta principios del siglo XX. Para evitar ser repetitivo he prescindido casi siempre de los calificativos honoríficos (el mayor, el más brillante) e insisto en la palabra «entre» de la frase anterior.

ANONIMOS egipcios, babilonios, chinos, mayas, indios y otros. Los escribas, monjes, astrónomos y pastores desconocidos que desarrollaron los conceptos y notaciones básicos de número y cifra.

PITAGORAS (hacia 540 a. C.). Griego. Junto con Tales fundó la matemática griega. El teorema de Pitágoras se atribuye a su escuela. Filósofo y místico de los números («Todo es número»).

PLATON (427?-347 a. C.). Griego. No fue propiamente un matemático, pero se le conoció como «el hacedor de matemáticos». Sus *Diálogos* y su

República están repletos de especulaciones matemáticas y referencias. Sobre la puerta de su academia escribió el lema: «No entre aquí nadie sin saber geometría».

EUCLIDES (hacia 300 a. C.). Griego. Autor de *Los Elementos*, sistematización de la geometría griega, que tuvieron un fuerte impacto sobre el pensamiento matemático durante milenios.

ARQUIMEDES (287-212 a. C.). Griego. Trabajó fundamentalmente en geometría, en concreto en el cálculo de áreas y volúmenes por el «método de exhaustación». Fue el mayor científico de la Antigüedad, autor de importantes descubrimientos en astronomía, hidrostática («*Eureka*, lo encontré») y mecánica.

APOLONIO (hacia 230 a. C.). Griego. Desarrolló nuevos aspectos de la geometría en sus *Secciones cónicas*. Precursor de la geometría analítica de Descartes. Astrónomo.

PTOLOMEO (907-160? d. C.). Griego alejandrino. Autor del *Almagesto* («el más grande»), que durante siglos fue libro de cabecera para la astronomía, la geografía y la matemática. Sistema ptolemaico.

DIOFANTO (hacia 250 d. C.). Griego alejandrino. Autor de la *Aritmética*, que trata de la teoría de los números y de las soluciones enteras de las ecuaciones.

AL-JUARIZMI (hacia 830). Árabe de Bagdad. Su *Al-jabr wa'l Muqabalah* fue el primer texto de álgebra elemental. Tuvo gran influencia en la transmisión de esta materia a Europa. Uno de los varios matemáticos árabes notables.

OMAR KHAYYAM (1050?-1123). Persa. Más conocido como autor del *Rubáiyát*. También fue un matemático notable cuyo trabajo en álgebra extendió el de Al-Juarizmi. Y **BHASKARA** (1114-1185). Indio. Principal algebrista del siglo XII.

GERONIMO CARDANO (1501-1576) y **NICCOLO TARTAGLIA** (1500-1557). Italianos. Descubridores de las fórmulas de las soluciones de las ecuaciones de tercer y cuarto grado. Estimularon la investigación en

álgebra.

GALILEO GALILEI (1564-1642). Italiano. Aunque no fuera matemático, su obra *Dos nuevas ciencias* contribuyó a forjar la alianza revolucionaria entre la matemática y la experimentación y los nuevos métodos para explorar las verdades de la naturaleza.

FRANÇOIS VIÈTE (1540-1603) Francés. Inventó una notación indispensable para el álgebra y, especialmente, las variables. Trabajó en álgebra y trigonometría. Los decimales de **SIMON STEVIN** (1548-1620), los logaritmos de **JOHN NAPIER** (1550-1617) y las elipses y estudios astronómicos de **JOHANNES KEPLER** (1571-1630) datan aproximadamente de la misma época.

RENÉ DESCARTES (1596-1650). Francés. Inventó, junto con Pierre de Fermat, la geometría analítica, uniendo los dos campos del álgebra y la geometría para sentar las bases de la moderna matemática. Filósofo, legó a la posteridad la duda cartesiana.

PIERRE DE FERMAT (1601-1665). Francés. Junto con Desearles inventó la geometría analítica. Trabajó en la teoría de los números, donde nos dejó su «último teorema», en cálculo y en teoría de la probabilidad.

BLAISE PASCAL (1623-1662). Francés. Uno de los fundadores de la teoría de la probabilidad. Apuesta de Pascal. Filósofo, místico religioso y estilista literario.

ISAAC NEWTON (1642-1727). Inglés. Inventó/descubrió el teorema del binomio, las series infinitas y el cálculo infinitesimal (simultáneamente con Leibniz); la matemática moderna data de su época. También fundó la física moderna, incluyendo las leyes del movimiento, la gravitación y la óptica en sus *Principios de la filosofía natural*.

GOTTFRIED WILHELM VON LEIBNIZ (1646-1716). Alemán. Inventó el cálculo (simultáneamente con Newton). Filósofo y lógico que anticipó muchos avances posteriores.

LOS BERNOUILLI, entre los que se incluyen **JAKOB** (1654-1705), **JOHANN** (1667-1748), **DANIEL** (1700-1782), y otros. Familia de matemáticos

suizos. Contribuciones importantes al cálculo de variaciones, la probabilidad y la física matemática.

LEONHARD EULER (1707-1783). Suizo. Prolífico en muchos campos de la matemática. Teoremas importantes en cálculo, ecuaciones diferenciales, teoría de los números, matemática aplicada y combinatoria.

JOSEPH LOUIS LAGRANGE (1736-1813). Francés. Teoría de los números, ecuaciones diferenciales, cálculo de variaciones, análisis, mecánica celeste y dinámica. Su *Mecánica analítica* es una bella obra de matemática pura.

PIERRE SIMON LAPLACE (1749-1827). Francés. Trabajos influyentes en análisis matemático y astronomía. Fundó la teoría de la probabilidad como una rama seria de la matemática. Contemporáneo del MARQUÉS DE CORDORCET (1743-1794) y sus aplicaciones sociales de la matemática y de los análisis de A. M. LEGENDRE (1752-1833).

NIKOLAI LOBACHEVSKI (1793-1856) y **JÁNOS BOLYAI** (1802-1860). Ruso y húngaro, respectivamente, descubridores (junto con Gauss) de la geometría no euclídea, en la que no es válido el postulado de las paralelas de Euclides.

KARL FRIEDRICH GAUSS (1777-1855). Alemán. Responsable de muchos resultados importantes en teoría de los números (*Investigaciones aritméticas*), teoría de las superficies, geometría no euclídea, física matemática y estadística (la curva gaussiana). Considerado junto a Newton y Arquímedes uno de los tres grandes matemáticos de todos los tiempos.

AUGUSTIN LOUIS CAUCHY (1789-1857). Francés. Teoremas importantes en variable compleja, otros trabajos en análisis y teoría de grupos. Uno de los iniciadores de los métodos formales y rigurosos en el cálculo.

ADOLPHE QUETELET (1796-1874). Belga. El primero en promover el uso de la probabilidad y los modelos estadísticos en la descripción de los fenómenos sociales, económicos y biológicos.

EVARISTE GALOIS (1811-1832). Francés. Su trabajo en la teoría de

ecuaciones determinó el estilo del álgebra abstracta moderna.

WILLIAM HAMILTON (1805-1865), **ARTHUR CAYLEY** (1821-1895) y **J. J. SILVESTER** (1814-1897). Irlandés e ingleses, respectivamente. Desarrollaron el álgebra abstracta como estudio de las operaciones, la forma y la estructura. Matrices. Cuaterniones.

BERNHARD RIEMANN (1826-1866). Alemán. Su trabajo más original sentó la base geométrica de la teoría de la relatividad general de Einstein y liberó la geometría de su dependencia de los conceptos físicos de longitud, anchura y altura.

GEORGE BOOLE (1815-1864). Inglés. Sus *Leyes del pensamiento* representaron un enfoque matemático del estudio de la lógica que tuvo una continuación importante en **GOTTLOB FREGE** (1848-1925) y **GIUSEPPE PEANO** (1858-1932).

JAMES CLERK MAXWELL (1831-1879). Escocés. No fue matemático, pero su desarrollo matemático de la teoría del campo electromagnético le hace merecedor de un lugar en esta lista, al igual que el trabajo matemático del norteamericano **JOSSIAH GIBBS** (1839-1903).

GEORG CANTOR (1845-1918). Alemán. Inventor/descubridor de la teoría de conjuntos, incluido el estudio de los conjuntos infinitos y los números transfinitos. Trabajó también en números irracionales y series infinitas.

FELIX KLEIN (1849-1925). Alemán. Unificó el estudio de la geometría examinando las propiedades invariantes bajo un grupo de transformaciones concreto. Maestro influyente.

JULES HENRI POINCARÉ (1854-1912). Francés. Contribuciones en muchos campos de la matemática, como la topología, las ecuaciones diferenciales, la física matemática y la probabilidad. También obras explicativas del método científico.

DAVID HILBERT (1862-1943). Alemán. Líder de una escuela formalista de matemática axiomática y autor de los *Fundamentos de geometría*, que suprimió los huecos lógicos de Euclides. Propuso una lista de problemas famosos no resueltos. Contribuyó al álgebra, la teoría de los espacios

abstractos (espacios de Hilbert) y al análisis (curvas que llenan el espacio).

G. H. HARDY (1877-1947), **J. E. LITTLEWOOD** (1855-1977) y **S. RAMANUJAN** (1887-1920). Ingleses e indio, respectivamente. Mucho trabajo original en análisis y teoría de los números, individualmente y en colaboración. Ramanujan se basaba generalmente en la pura intuición. También otras obras más generales.

BERTRAND RUSSELL (1872-1970). Inglés. Autor de los *Principia Mathematica* (con ALFRED NORTH WHITEHEAD), donde se deduce toda la matemática únicamente (casi) a partir de la lógica. Paradoja de Russell. Filósofo y divulgador. (Aunque éste no sea uno de los méritos que le han hecho famoso, Russell fue uno de mis héroes de juventud en el instituto y, como incluyó en su autobiografía una carta que le escribí cuando iba a la universidad, sigue siéndolo hoy en gran medida).

KURT GÖDEL (1906-1978). Austríaco/norteamericano. Su teorema de incompletitud estableció la existencia de proposiciones indecidibles en todas las formalizaciones de la matemática. Otras obras muy originales sobre consistencia, intuicionismo y recurrencia.

ALBERT EINSTEIN (1879-1955). Alemán/norteamericano. Aunque no fue específicamente un matemático, su trabajo revolucionario en relatividad general (por no hablar de todas sus demás aportaciones) le hace merecedor de un lugar en esta lista.

JOHN VON NEWMANN (1903-1957). Húngaro/norteamericano. Aportaciones decisivas en la fundamentación de la matemática, la física matemática (especialmente la teoría cuántica), el análisis y el álgebra abstracta. Inventó la teoría de juegos. Trabajos importantes en ordenadores y autómatas.

La lista omite necesariamente muchos grandes matemáticos del pasado. Tampoco contiene *ningún* matemático actual, y por varias razones: son demasiado numerosos para citarlos (contando sólo los asociados a las

principales organizaciones profesionales norteamericanas, hay más de 30.000 nombres, y centenares de revistas de investigación). El veredicto sobre su obra todavía no está claro; sus resultados son a veces demasiado especializados para resumirlos ni siquiera superficialmente, y hay muchos autores y campos periféricos que quizás habría que considerar también — Benoît Mandelbrot y sus fractales, estadísticos como Ronald Fisher y Karl Pearson, informáticos como Alan Turing, Marvin Minsky y Donald Knuth, A. N. Kolmogorov y su trabajo sobre probabilidad y complejidad, Claude Shannon sobre teoría de la información, Kenneth Arrow sobre las funciones de elección social, y muchos más.

Un examen superficial de la lista revela también que no hay mujeres, negros ni orientales. La tendencia reciente a acogerlas en el Parnaso literario no ha afectado al panteón matemático. No se trata, desde luego, de que estos grupos carezcan de un talento matemático al más alto nivel. Desde Hypatia, una mujer alejandrina que escribió comentarios matemáticos y cayó víctima de una turba de cristianos fanáticos en el año 415, a la hija de Lord Byron, Ada Lovelace, que escribía programas para la máquina analítica de Charles Babbage (un elaborado ordenador mecánico), y a Emmy Noether, eminente algebrista que fue expulsada por los nazis y enseñó en el Bryn Mawr College en los años treinta, las mujeres han tenido que luchar por el requisito mínimo de Virginia Woolf de un lugar propio. Esta situación ha mejorado un poco y en los últimos años una minoría importante de los doctorados en matemáticas han sido para mujeres. (Nuevas investigaciones indican que la mayoría de mujeres matemáticas han contado con el apoyo familiar: muchas son hijas de matemáticos, en su juventud han tenido contacto con verdadera matemática y con modelos del papel femenino en este campo).

En cuanto a los orientales, los negros y otros grupos étnicos, no hay que buscar demasiado lejos su perspicacia matemática para disipar el eurocentrismo histórico de la lista. Considérense, por citar unos pocos ejemplos al azar, la larga historia de la matemática china (mal conocida en el oeste: el triángulo de Pascal fue descubierto 300 años antes de Pascal, por ejemplo), los logros matemáticos y cuasimatemáticos de muchas culturas «primitivas», las realizaciones de los egipcios, indios, africanos y árabes, y el trabajo de investigación de los matemáticos japoneses de hoy. Cualesquiera

que hayan sido las restricciones y limitaciones del pasado, hoy la matemática se estudia en todos los países y continentes del mundo. No obstante, el tópico de la universalidad de la matemática probablemente no será plenamente apreciado hasta que se dé la correspondiente universalidad de los matemáticos.

Lecturas recomendadas

En su mayoría, los siguientes libros son accesibles al público en general. Contienen pocas ecuaciones y se centran en las ideas matemáticas, y no en cálculos pesados ni en demostraciones rigurosas. El lector que tenga ya unos conocimientos matemáticos sólidos probablemente sabrá ya dónde profundizar más en cada tema.

Abbot, Edwin A., *Flatland*, Dover Publications, 1952.

Albers, Donald J., y G. L. Alexanderson, eds., *Mathematical People*, Birkhauser, 1985.

Beckmann, Petr, *A History of Pi*, St. Martin's Press, 1971.

Bell, Eric Temple, *Mathematics: Queen and Servant of Science*, Mathematical Association of America, 1987.

Benacerraf, Paul, y Hilary Putnam, eds., *Philosophy of Mathematics*, Prentice-Hall, 1964.

Boyer, Carl, *The History of Calculus and Its Conceptual Development*, Dover Publications, 1959.

—, *A History of Mathematics*, John Wiley & Sons, 1968. Traducción española: *Historia de las matemáticas*, Alianza Editorial, Madrid, 1987.

Chinn, W. G., y N. E. Steenrod, *First Concepts of Topology*, Mathematical Association of America, 1966. Hay traducción española: *Primeros conceptos de topología*, Alhambra, Madrid, 1975.

- COMAP (Consortium for Mathematics and Its Applications), *For All Practical Purposes: Introduction to Contemporary Mathematics*, W. H. Freeman, 1988.
- Courant, Richard, y Herbert Robbins, *What Is Mathematics?*, Oxford University Press, 1948.
- Daintih, John, y R. D. Nelson, eds., *Dictionary of Mathematics*, Penguin Books, 1989.
- Davis, Philip J., *The Lore of Large Numbers*, Mathematical Association of America, 1961.
- Davis, Philip J., y Reuben Hersh, *The Mathematical Experience*, Houghton Mifflin, 1981. Hay traducción española: *Experiencia matemática*, Labor, Barcelona, 1989.
- Eves, H., y C. V. Newson, *An Introduction to the Foundations and Fundamental Concepts of Mathematics*, Holt, Rinehart and Winston, 1965.
- Gardner, Martin, *Aha! Insigth*, W. H. Freeman, 1978. Hay traducción española: *¡Ajá! Inspiración*, Labor, Barcelona, 1981.
- , *Aha! Gotcha*, W. H. Freeman, 1982. Hay traducción española: *Ajá*, Labor, Barcelona, 1985.
- , *Wheels, Life and Other Mathematical Amusements*, W. H. Freeman, 1983. Hay traducción española: *Ruedas, vida y otras diversiones matemáticas*, Labor, Barcelona, 1988.
- , *Penrose Tiles to Trapdoor Ciphers*, W. H. Freeman, 1989. Hay traducción española: *Mosaicos de Penrose y escotillas cifradas*, Labor, Barcelona, 1990.
- Gleick, James, *Chaos: Making a New Science*, Viking Press, 1987. Existe traducción española: *Caos*, Seix Barral, Barcelona, 1988.
- Guillen, Michael, *Bridges to Infinity*, Jeremy P. Tarcher, 1983.
- [Guzmán, Miguel de, *Aventuras matemáticas*, Labor, Barcelona, 1988].
- Hardy, G. H., *A Mathematician's Apology*, Cambridge University Press,

1967.

Heath, T. L., *A History of Greek Mathematics*, Clarendon Press, 1967.

Hofstadter, Douglas, *Gödel, Escher, Bach*, Basic Books, 1980. Existe traducción española: *Gödel, Escher y Bach*, Tusquets Editores (Metatemas 14), Barcelona, 1992 (4.^a edición).

—, *Metamagical Themas*, Basic Books, 1985.

Ifrah, Georges, *From One to Zero: A Universal History of Numbers*, Viking Press, 1985. Existe traducción española: *Las cifras: historia de una gran invención*, Alianza Editorial, Madrid, 1988.

Kac, Mark, y Stanislaw Ulam, *Mathematics and Logic*, Frederick Praeger, 1968.

Kline, Morris, *Mathematical Thought from Ancient to Modern Times*, Oxford University Press, 1972.

—, *Mathematics: An Introduction to Its Spirit and Use*, W. H. Freeman, 1977.

—, *Mathematics for the Non-mathematician*, Dover Publications, 1985. Existen traducciones españolas de otras dos obras: *El fracaso de las matemáticas modernas*, Siglo XXI, Madrid, 1986, y *Matemáticas. La pérdida de la certidumbre*, Siglo XXI, Madrid, 1985.

Littlewood, J. E. (Bela Bollobas, editor), *Littlewood's Miscellany*, Cambridge University Press, 1972.

Mandelbrot, Benoît, *The Fractal Geometry of Nature*, W. H. Freeman, 1977. Hay traducción española: *La geometría fractal de la naturaleza*, Tusquets Editores, Barcelona, en prensa.

Mason, John, Leo Burton y Kaye Istacey, *Thinking Mathematically*, Alison Wesley, 1987. Hay traducción española: *Pensar matemáticamente*, MEC/Labor, Barcelona, 1992.

[Mataix, Mariano, *Cajón de sastre matemático*, Marcombo, Barcelona, 1993].

Moore, David S., *Statistics: Concepts and Controversies*, W. H. Freeman,

1979.

Niven, Ivan, *The Mathematics of Choice*, Mathematical Association of America, 1965.

Ore, Oystein, *Graphs and Their Uses*, Mathematical Association of America, 1963.

Packel, Edward, *The Mathematics of Games and Gambling*, Mathematical Association of America, 1981.

Pappas, Theoni, *The Joy of Mathematics*, Wide World Publishing/Tetra, 1989.

Paulos, J. A., *Mathematics and Humor*, University of Chicago Press, 1980.

—, *I Think, Therefore I Laugh*, Columbia University Press, 1985. Existe traducción española: *Pienso, luego río*, Cátedra, Madrid, 1988.

—, *Innumeracy: Mathematical Illiteracy and Its Consequences*, Farrar, Straus & Giroux, 1989. Traducción española: *El hombre anumérico*, Tusquets Editores (Metatemas 20), Barcelona, 1990 (2.^a edición).

Peterson, Ivars, *The Mathematical Tourist*, W. H. Freeman, 1988. Existe traducción española: *El turista matemático*, Alianza Editorial, Madrid, 1992.

—, *Islands of Truth*, W. H. Freeman, 1990.

Polya, George, *How to Solve It*, Princeton University Press, 1945.

Poundstone, William, *The Recursive Universe*, William Morrow, 1985.

[Rodríguez Vidal, R., *Enjambres matemáticos*, Reverté, Barcelona, 1992].

Rucker, Rudy, *Mind Tools*, Houghton Mifflin, 1987.

Russell, Bertrand, *The Principles of Mathematics*, Cambridge University Press, 1903. Existe traducción española: *Los principios de la matemática*, Espasa-Calpe, Madrid, 1983.

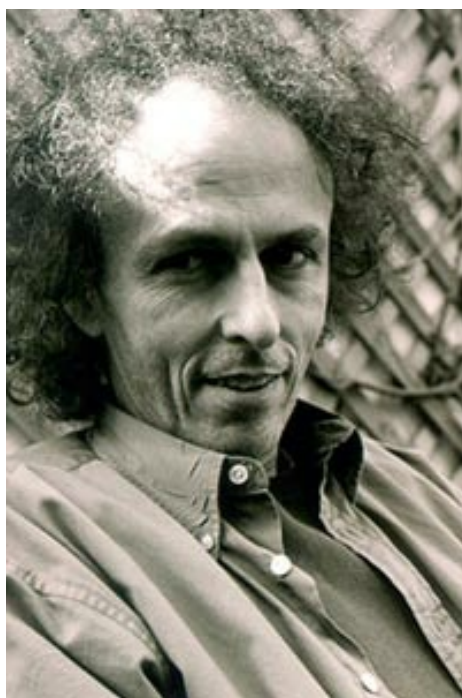
Salmon, Weslwy, *Space, Time, and Motion*, Dickenson, 1975.

Steen, Lynn Arthur, ed., *Mathematics Today*, Springer-Verlag, 1978.

Weaver, Warren, *Lady Luck*, Dover Publication, 1982.

Weizenbaum, Joseph, *Computer Power and Human Reason*, W. H. Freeman, 1976. Existe traducción española: *La frontera entre el ordenador y la mente*, Pirámide, Madrid, 1977.

Wilder, Raymond, *The Foundations of Mathematics*, John Wiley & Sons, 1965.



JOHN ALLEN PAULOS. Denver (EE. UU.), 1945. Doctor en matemáticas por la Universidad de Wisconsin y profesor de esta materia en la Temple University de Filadelfia. Además de escritor de éxito, es un afamado conferenciante, comentarista y respetado columnista sobre disciplinas como la filosofía de la ciencia, la lógica y las matemáticas, así como sobre las hilarantes aberraciones que la ignorancia matemática suele generar. Ha comentado asimismo decenas de libros en publicaciones como *The New York Review of Books* o *The London Review of Books*. En 2002 recibió el University Creativity Award y, en 2003, el American Association for the Advancement of Science Award, por su contribución a la divulgación de la ciencia.

Notas

[1] *El hombre anumérico*, Tusquets Editores (Metatemas 20), Barcelona, 1989. (N. del T.) <<

[2] Uso coloquial y descuidado de la partícula *don't*, en lugar de *doesn't*, para el singular de la tercera persona. (N. del T.) <<

[3] Literalmente: «No paso porque el peligro se extiende por doquier los locos de la velocidad proliferan incontrolados». (*N. del T.*) <<

[4] Food and Drug Administration, equivalente en Estados Unidos a nuestro ministerio de sanidad y consumo. (*N. del T.*) <<

[5] Broma intraducible: *difi-q* se pronuncia como las primeras silabas *diff. eq.* que abrevian *differential equations* y a la vez suena como *difficult*, difícil. (N. del T.) <<

[6] En junio de 1993, el matemático británico Andrew Wiles anunció públicamente la demostración final, superando un mito tras 356 años. (*N. del T.*) <<

[*] La demostración de Wiles que publicó en 1993 tenía un error. Después de dos años de duro trabajo, pudo corregir el error y publicar en 1995 la demostración final del teorema. (*N. del E. D.*) <<

^[7] Massachussets Institute of Technology. (*N. del T.*) <<

[8] Se refiere a la piratería informática, no en el sentido de copiar, usar o vender *software* sin la debida licencia, sino en el del que se introduce en un ordenador ajeno a través de una red informática. (*N. del T.*) <<

[9] Literalmente: «*Que te chodan, Jarlie*». Sería el juego de palabras empleado en la expresión «patiabierto y boquidifuso». (*N. del T.*) <<

[10] Mezcla cómica de dos proverbios. Equivaldría a «Poner las cartas sobre las íes» o «Más vale pájaro en mano que te sacarán los ojos». (*N. del T.*) <<

[11] Literalmente: «No apruebo a Sally a no ser que parezca temerario». (*N. del T.*) <<

[12] La broma es intraducible. En inglés *pi*, el número, se pronuncia «pai», igual que *pie*, pastel, cuya forma circular permite el chiste. (*N. del T.*) <<

[13] Chiste intraducible: la fonética norteamericana de *solids* se puede transformar fácilmente en la de *salads*. Cambiando el personaje antiguo «Platón» por «César», el chiste está servido. (*N. del T.*) <<

[14] En España lo fueron, hace unos 20 años, los envases de un conocido refresco. (*N. del T.*) <<

[15] Chiste intraducible: *featherbed*, «subvención excesiva», pero también «plumón». (*N. del T.*) <<

[16] Syene corresponde a la actual Asuán. (*N. del T.*) <<

[17] En pronunciación inglesa: *Sin X* = sainex; *Somincx* = somainex;
somnífero = *somnijerous*. (*N. del T.*) <<

[18] Intraducible: *Four* = 4; *Fourier* se pronuncia «four-ye» igual que 4 «ye», y *fortran* se pronuncia «for-tran», igual que 4 «tran». (N. del T.) <<